

Claude Servin

Aide-mémoire de **RÉSEaux ET TÉLÉCOMS**



- ▶ **Notions de base sur les réseaux**
- ▶ **Principes de l'architecture TCP/IP**
- ▶ **Réseaux locaux et réseaux d'opérateurs**
- ▶ **Qualité de service et sécurité**



DUNOD

VISIT...

LANZAROTE
Caliente.COM

Claude Servin

Aide-mémoire des
RÉSEAUX ET TÉLÉCOMS

DUNOD

photo de couverture :
© Vincent TT - fotolia.com

Le pictogramme qui figure ci-contre mérite une explication. Son objet est d'alerter le lecteur sur la menace que représente pour l'avenir de l'écrit, particulièrement dans le domaine de l'édition technique et universitaire, le développement massif du photocopillage.

Le Code de la propriété intellectuelle du 1^{er} juillet 1992 interdit en effet expressément la photocopie à usage collectif sans autorisation des ayants droit. Or, cette pratique s'est généralisée dans les établissements

d'enseignement supérieur, provoquant une baisse brutale des achats de livres et de revues, au point que la possibilité même pour

les auteurs de créer des œuvres nouvelles et de les faire éditer correctement est aujourd'hui menacée.

Nous rappelons donc que toute reproduction, partielle ou totale, de la présente publication est interdite sans autorisation de l'auteur, de son éditeur ou du

Centre français d'exploitation du droit de copie (CFC, 20, rue des Grands-Augustins, 75006 Paris).



© Dunod, Paris, 2004, 2009, 2012

978-2-10-058681-3

Le Code de la propriété intellectuelle n'autorisant, aux termes de l'article L. 122-5, 2^o et 3^o a), d'une part, que les « copies ou reproductions strictement réservées à l'usage privé du copiste et non destinées à une utilisation collective » et, d'autre part, que les analyses et les courtes citations dans un but d'exemple et d'illustration, « toute représentation ou reproduction intégrale ou partielle faite sans le consentement de l'auteur ou de ses ayants droit ou ayants cause est illicite » (art. L. 122-4).

Cette représentation ou reproduction, par quelque procédé que ce soit, constituerait donc une contrefaçon sanctionnée par les articles L. 335-2 et suivants du Code de la propriété intellectuelle.

Table des matières

Avant-propos	XI
Partie 1 Notions de base sur les réseaux	1
1. Qu'est-ce qu'un réseau ?	3
2. Le modèle OSI	8
2.1 Description du modèle de référence	8
2.2 Principes de base d'une architecture en couche	10
3. L'architecture TCP/IP	12
3.1 Origine	12
3.2 Principe architectural	13
3.3 Description générale de l'environnement TCP/IP	14
3.4 Conclusion	16
Partie 2 Éléments physiques de la liaison de données	17
4. Les supports et leur limitation	19
4.1 Les supports cuivre	19
4.2 Les supports guidés	23
4.3 Les supports non guidés	40
5. Les modes de transmission	45
5.1 L'organisation des échanges	45
5.2 Transmissions bande de base et large bande	50
5.3 La transmission large bande	56
6. Le multiplexage	59
6.1 Le multiplexage spatial	60
6.2 Le multiplexage temporel	62

6.3	Le multiplexage inverse	63
6.4	Conclusion	64
Partie 3 Les protocoles de liaison		65
7. Les fonctions élémentaires		67
7.1	Notion de protocole	67
7.2	La délimitation des données	67
7.3	Le contrôle d'erreur	69
8. Exemples de protocole de liaison		73
8.1	SLIP (Serial Line Internet Protocol)	73
8.2	HDLC, High Level Data Link Control	74
8.3	PPP (Point to Point Protocol)	88
8.4	Conclusion	97
Partie 4 Le niveau réseau		99
9. Le concept de réseau à commutation		101
9.1	Définitions	101
9.2	Les réseaux à commutation	103
9.3	Performances des réseaux à commutation	107
10. Les réseaux à commutation de paquets		109
10.1	Du mode datagramme au mode connecté	109
11. Les techniques réseau		114
11.1	La notion d'adressage	114
11.2	La segmentation et le réassemblage	118
11.3	Le contrôle de congestion	119
11.4	L'acheminement	121
12. Le réseau IP		126
12.1	L'adressage dans IP	126
12.2	L'adressage dans le réseau logique	127
12.3	Les techniques d'adressage dans un réseau IP	128
12.4	La structure du datagramme IP	136
12.5	Le contrôle de la fragmentation sous IP	140
13. D'IPv4 à IPv6		142
13.1	Les lacunes d'IPv4	142
13.2	L'adressage dans IPv6	143

13.3	Le datagramme IPv6	146
13.4	Conclusion	148
Partie 5	Les protocoles de transport : TCP et UDP	149
14.	Les mécanismes de base de TCP	151
14.1	La notion de connexion de transport	151
14.2	Etablissement de la connexion de transport	154
14.3	Le mécanisme contrôle de l'échange	155
14.4	L'option d'estampille horaire	160
15.	Le segment TCP et les mécanismes associés	161
15.1	La structure du segment TCP	161
15.2	Le contrôle d'erreur	162
15.3	La taille des segments	163
15.4	Le TCP et les réseaux à haut débit	164
16.	TCP/UDP et le multimédia, le contrôle de flux et de congestion	165
16.1	Définitions	165
16.2	Le contrôle de flux	166
16.3	Le contrôle de la congestion	166
16.4	UDP dans IPv4	171
16.5	UDP dans IPv6	172
16.6	Conclusion	172
Partie 6	TCP/IP utilitaires et applications	175
17.	Les utilitaires de la couche réseau	177
17.1	Le protocole ICMP	177
17.2	La résolution d'adresses	181
17.3	Les utilitaires de configuration dans IPv4	185
17.4	L'auto-configuration dans IPv6	187
17.5	IP et la mobilité	188
18.	Les applications de l'environnement TCP	192
18.1	Notions d'annuaire	192
18.2	Le transfert de fichiers	198
18.3	L'émulation de terminal (Telnet)	200
18.4	La messagerie électronique	201
18.5	Les notions de middleware	204

Partie 7 Les réseaux locaux	211
19. Le réseau local	213
19.1 Les constituants d'un réseau local	213
19.2 Les réseaux locaux et la normalisation	214
19.3 La couche physique	214
19.4 La sous-couche MAC	215
19.5 La couche liaison (LLC)	218
20. Les réseaux Ethernet	220
20.1 Présentation	220
20.2 Caractéristiques des réseaux Ethernet	222
20.3 Les différentes versions d'Ethernet	224
21. La commutation dans les LAN – Les réseaux virtuels ou VLAN	230
21.1 Principe de la commutation dans les LAN	230
21.2 Ethernet <i>full duplex</i>	231
21.3 Principes généraux des VLAN	232
21.4 L'identification des VLAN (802.1Q)	233
22. L'Ethernet sans fil	236
22.1 Généralités	236
22.2 La problématique de l'accès aux réseaux sans fil	236
22.3 L'architecture générale des réseaux sans fil	237
22.4 Les réseaux 802.11	238
22.5 Conclusion	248
Partie 8 Les réseaux d'opérateur	249
23. Structure et protocoles	251
23.1 Architecture générale	251
23.2 Structure générale d'un réseau	252
23.3 Le plan de transmission	253
23.4 Le plan de service	254
24. MPLS, Multiprotocol Label Switching	270
24.1 Principe	270
24.2 Le réseau MPLS	271
24.3 Les VPN MPLS	275

25. L'accès aux réseaux, la boucle locale	278
25.1 Définition	278
25.2 Organisation de la distribution des accès	278
25.3 Les accès haut débit	279
26. Ethernet dans les MAN et WAN	284
26.1 Les réseaux sans coupure	284
26.2 Ethernet à grande distance (CGE, Carrier Grade Ethernet)	285
26.3 Modèle architectural	287
27. Sécurisation des accès	289
27.1 Conclusion	291
Partie 9 Interconnexion des réseaux et la qualité de service	293
28. L'interconnexion des réseaux	295
28.1 Définition	295
28.2 Les problèmes liés à l'interconnexion	295
28.3 L'encapsulation ou <i>tunneling</i>	296
29. Les éléments d'interconnexion (relais)	298
29.1 Définitions	298
29.2 Les routeurs	299
30. Les techniques de routage	301
30.1 Généralités	301
30.2 Le routage dans le réseau IP	302
31. Le routage <i>multicast</i>	313
31.1 Introduction au <i>multicast</i>	313
31.2 Le protocole local IGMP (RFC 2236)	314
31.3 Les protocoles de routage <i>multicast</i>	314
31.4 Internet et le <i>multicast</i>	317
Partie 10 La téléphonie sur IP	319
32. Principes généraux de la téléphonie	321
32.1 Introduction	321
32.2 De l'analogique à la ToIP	322
32.3 Notions d'autocommutateurs privés	329

33. La téléphonie sur IP	334
33.1 Généralités	334
33.2 La téléphonie, une application parmi d'autres ?	335
34. L'architecture logique et la signalisation	340
34.1 L'architecture H.323 de l'UIT-T	341
34.2 Le protocole SIP de l'IETF (RFC 3261)	344
34.3 Signalisation, la synthèse	349
35. Mise en œuvre de la ToIP	350
35.1 L'architecture générale	350
35.2 La qualité de service	357
35.3 Conclusion	359
Partie 11 La sécurité des systèmes d'informations	361
36. La sécurité des données	363
36.1 Généralités	363
36.2 La protection des données	364
37. La sécurisation des échanges	370
37.1 L'usurpation d'identité	370
37.2 La sécurité et le protocole de transmission	371
38. La sécurisation du réseau	375
38.1 Les menaces	375
38.2 La protection de l'intranet	376
38.3 Conclusion	383
Annexes	385
Index	395

AVANT-PROPOS

Avec Internet, les réseaux sont sortis du domaine de l'entreprise et ont envahi le domicile. Du routeur ADSL à l'imprimante WiFi, le « *home réseaux* » met en œuvre des techniques similaires à celles des réseaux d'entreprise. Aussi, l'étudiant en informatique, le technicien réseau en entreprise et le simple particulier curieux doivent tous posséder une connaissance des réseaux allant de la simple connectivité aux aspects de sécurité. Le but de cet ouvrage est donc d'aborder succinctement mais avec précision toutes les connaissances nécessaires à la bonne compréhension des techniques réseaux.

La vocation de la collection aide-mémoire étant d'apporter rapidement la réponse à une question, cet ouvrage a été divisé en 38 thèmes regroupés en 11 parties ce qui lui confère exhaustivité, précision et synthèse.

La première partie « Notions de base sur les réseaux » introduit la notion d'architecture notamment TCP/IP. La seconde décrit les éléments d'une liaison de données, notamment les supports, en s'attardant sur les problèmes de câblage de paires torsadées. La troisième rappelle les principes des protocoles et présente HDLC, PPP... La quatrième aborde la problématique des réseaux avec le concept d'adressage, d'acheminement et de contrôle. Les parties 5 et 6 sont consacrées aux protocoles de l'environnement TCP/IP. Les parties 7, 8, 9 sont dédiées aux réseaux d'opérateur, à la mise en œuvre et à la qualité de service de leurs réseaux. La dixième partie se consacre à l'étude de la téléphonie et tout particulièrement à la téléphonie sur IP tandis que la onzième et dernière aborde la sécurité.

Ainsi, le professionnel et l'étudiant curieux trouveront dans les pages qui suivent le rappel de toutes les connaissances fondamentales pour comprendre, mettre en œuvre et entretenir avec discernement un réseau.



Le lecteur trouvera sur le site de l'éditeur www.dunod.com *un lexique de toutes les abréviations et acronymes utilisés dans cet ouvrage.*

Note de l'éditeur

Pour aller plus loin, l'ouvrage du même auteur, *Réseaux et télécoms* (Dunod, 2009, coll. Science Sup, 1008 pages), propose un cours détaillé avec 129 exercices et problèmes à résoudre.

1

Notions de base sur les réseaux

1

Qu'est-ce qu'un réseau ?

En informatique le terme réseau recouvre un ensemble de moyens technologiques et logiciels mis en œuvre pour permettre l'échange de données entre ordinateurs. Cependant, en fonction des distances séparant les locuteurs, les techniques utilisées diffèrent. Ainsi a-t-on défini une classification essentiellement basée sur le critère distance (figure 1.1).

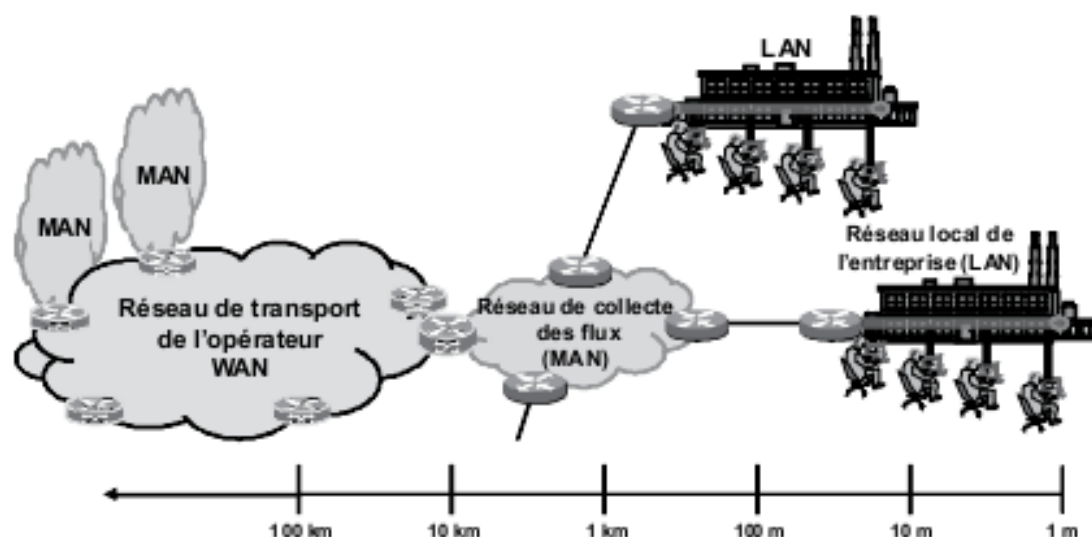


Figure 1.1 Relation entre les différents réseaux.

Cette classification traditionnelle correspond à un ensemble de contraintes que le concepteur devra prendre en compte lors de la réalisation de son réseau :

- **LAN** (Local Area Network), la notion de réseau local englobe un ensemble de techniques allant de celles nécessaires à la communication de plusieurs centaines de machines d'un même établissement d'une

entreprise à celles beaucoup plus simples mises en œuvre par un particulier pour relier son ordinateur et son imprimante à sa connexion Internet. Ces deux approches peuvent utiliser des techniques filaires ou radio.

- ▶ **MAN** (Metropolitan Area Network), d'une étendue de l'ordre d'une centaine de kilomètres, les MAN sont généralement utilisés pour fédérer les réseaux locaux ou assurer la desserte informatique de circonscriptions géographiques importantes (réseaux de campus).
- ▶ **WAN** (Wide Area Network), ces réseaux assurent l'acheminement des informations sur de grandes distances. Lorsque ces réseaux appartiennent à des opérateurs, les services sont offerts à des abonnés contre une redevance. Le réseau Internet n'a lui aucune existence propre, il est constitué d'un ensemble de réseaux d'opérateurs interconnectés entre eux (réseaux de réseaux).

De la détermination d'une route dans un réseau longue distance (WAN) à la localisation de la machine finale sur le réseau local (LAN), les techniques diffèrent et leur complexité aussi, mais le fondement des échanges reste la communication entre deux machines généralement appelées **nœuds du réseau**, que ceux-ci soient des **nœuds intermédiaires** dans un réseau WAN ou les machines terminales d'un réseau local. Cette relation directe entre deux nœuds s'appelle **une liaison point à point** (figure 1.2).



Figure 1.2 Liaison point à point.

Une liaison point à point met en relation 2 systèmes informatiques, que ceux-ci soient des ordinateurs terminaux ou des nœuds intermédiaires d'un réseau, et met en œuvre divers éléments et un ensemble de règles d'échange désigné sous le terme de **protocole**.

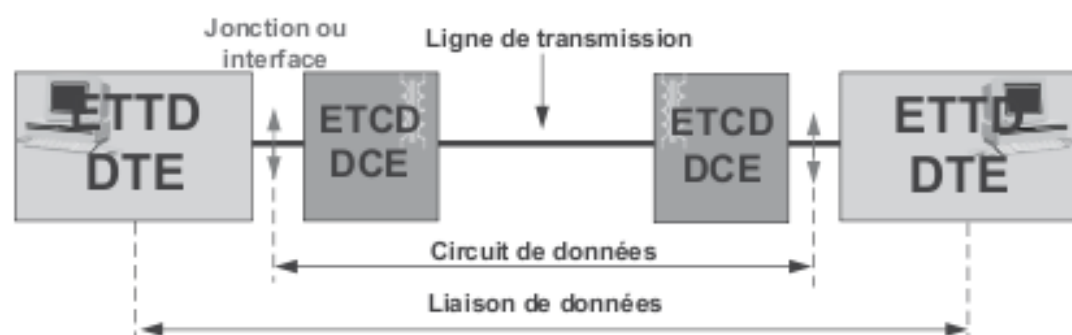


Figure 1.3 Constituant de base d'une liaison de données.

La figure 1.3 distingue :

- ▶ Les équipements terminaux (*End system*) ou **ETTD** (Équipement terminal de traitement de données), appelés aussi **DTE** (*Data Terminal Equipment*), ce sont soit des calculateurs d'extrémité soit les nœuds intermédiaires d'un réseau pris deux à deux. Ces machines sont dotées de circuits particuliers pour contrôler les communications (contrôleur de transmission). L'ETTD réalise la fonction de contrôle du dialogue.
- ▶ Des équipements d'adaptation ou **ETCD** (Équipement terminal de circuit de données), ou **DCE** (*Data Circuit Equipment*), réalisent l'adaptation entre les calculateurs d'extrémité et le support de transmission. Ces éléments remplissent essentiellement des fonctions électroniques ; ils transforment les données à transmettre en signaux adaptés aux caractéristiques du support de transmission. Ils modifient la nature du signal, mais pas sa signification.
- ▶ La jonction, interface entre ETTD (DTE) et ETCD (DCE), permet à l'ETTD de gérer l'ETCD afin d'assurer un déroulement correct des communications (établissement du circuit, initialisation de la transmission, échange de données et libération du circuit).
- ▶ Enfin, le support ou ligne de transmission, élément passif essentiel à la liaison et qui conditionne fortement les performances d'un système de transmission.

La liaison point à point relie physiquement deux nœuds, les problèmes à résoudre sont donc des problèmes de connectivité, d'adaptation au

support, de codage et décodage des informations. Alors que dans un réseau WAN, l'essentiel est la détermination d'une route afin d'assurer le transfert des informations entre l'installation source et l'installation destinatrice où il sera alors nécessaire d'identifier la machine cible sur le LAN local. La figure 1.4 illustre les différences essentielles d'acheminement entre les WAN et les LAN.

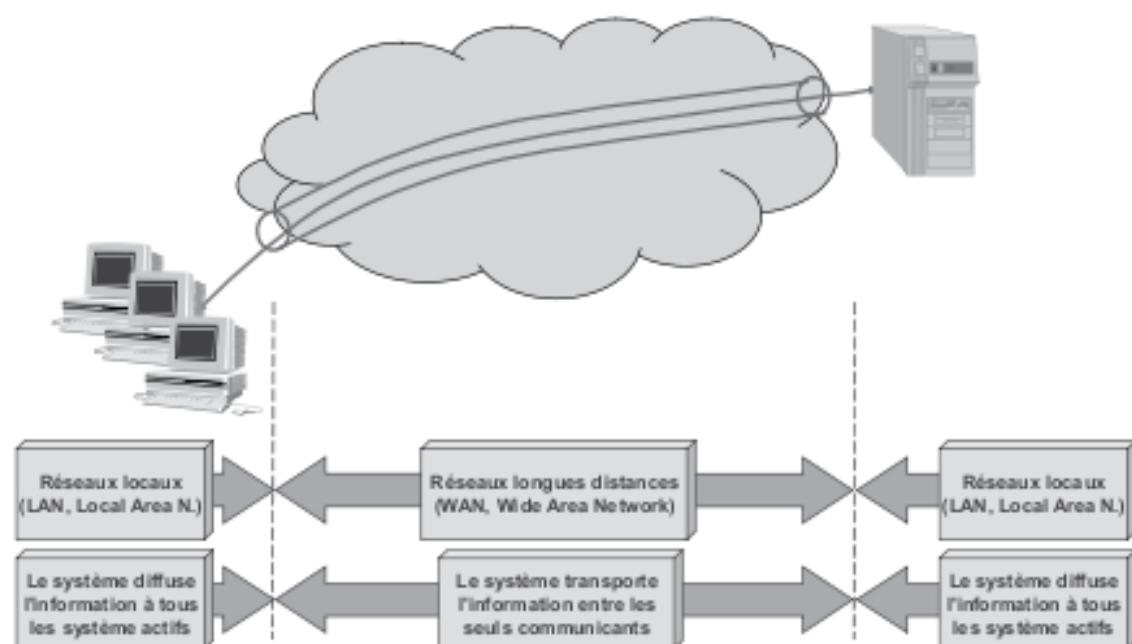


Figure 1-4 LAN et WAN, les différences fondamentales.

La différence de nature des problèmes évoqués succinctement ci-dessus, la complexité croissante des besoins de communication et la diversité des solutions adoptées ont très vite fait apparaître la nécessité de définir un cadre de développement complet qui « normalise » les solutions, ce modèle porte le nom d'**architecture protocolaire de réseau**.

Historiquement, chaque grand constructeur avait défini la sienne : SNA (*System Network Architecture*) pour IBM, DSA (*Distributed System Architecture*) pour BULL... Ces architectures propriétaires incompatibles entre elles ne permettaient pas l'interopérabilité des systèmes. Aussi, convenait-il

de définir des techniques de mise en relation en spécifiant une architecture normalisée. C'est ce qu'entreprend l'ISO (*International Standardization Organization*)¹ en définissant une architecture de communication normalisée, couramment appelée modèle de référence ou modèle **OSI** (*Open System Interconnection*).

1 C'est une habitude franco-française de traduire le terme ISO en International Standardization Organization ; en fait le nom officiel de l'ISO est : « International Organization for Standardization » et c'est parce que le nom de l'Organisation internationale de normalisation donnerait lieu à des abréviations différentes selon les langues (« IOS » en anglais et « OIN » en français), qu'il a été décidé d'emblée d'adopter un mot dérivé du grec *isos*, signifiant « égal ». La forme abrégée du nom de l'organisation est par conséquent toujours ISO (extrait du site officiel de l'ISO - www.iso.org). Il n'y a donc pas de traduction réelle de ce terme, ISO n'est pas un acronyme.

2

Le modèle OSI

2.1 Description du modèle de référence

Après de nombreux débats, le modèle de référence a identifié 7 grandes fonctionnalités définies en sept couches. En effet, pour réaliser une communication à travers un ou plusieurs systèmes intermédiaires (relais), il faut (figure 2.1) :

- ▶ relier les systèmes par un lien physique (couche ou niveau PHYSIQUE) ;
- ▶ contrôler qu'une liaison est correctement établie sur ce lien (couche ou niveau LIAISON) ;
- ▶ assurer à travers le relais (réseau) l'acheminement des données et la délivrance au bon destinataire (couche ou niveau RESEAU) ;
- ▶ contrôler, avant de délivrer les données à l'application que le transport s'est réalisé correctement de bout en bout (couche ou niveau TRANSPORT) ;
- ▶ organiser le dialogue entre toutes les applications, en gérant des sessions d'échange (couche ou niveau SESSION) ;
- ▶ traduire les données selon une syntaxe d'échange compréhensible par les deux entités d'application (couche ou niveau PRÉSENTATION) ;
- ▶ fournir à l'application utilisateur tous les mécanismes nécessaires pour masquer à celle-ci les contraintes de transmission (couche ou niveau APPLICATION).

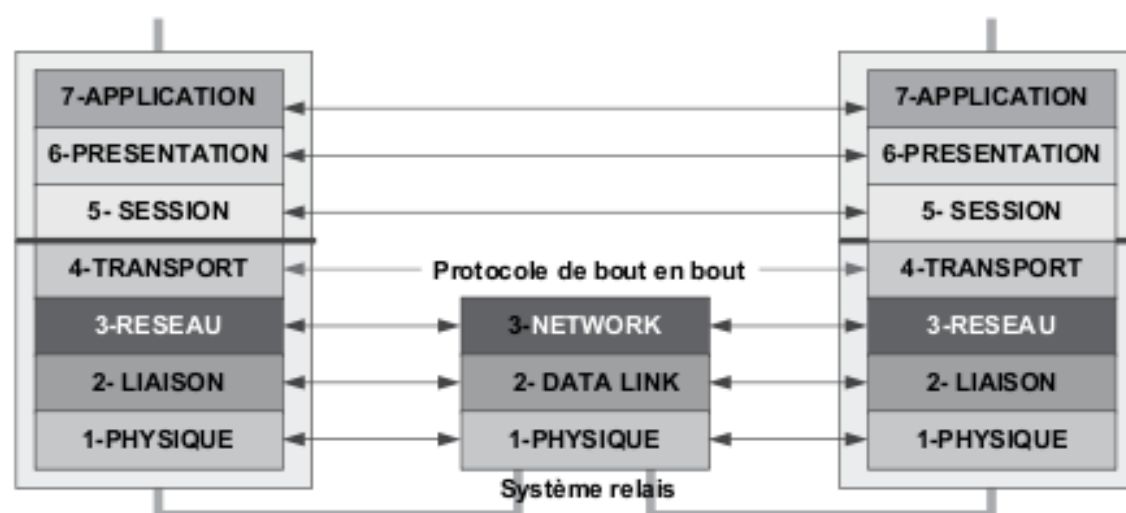


Figure 2.1 Le modèle de référence.

Le tableau 2.1 présente une synthèse des fonctionnalités de chacune des couches composant le modèle.

Tableau 2.1 Synthèse des fonctionnalités de chaque couche.

COUCHES	FONCTIONS
NIVEAU 1 Couche Physique Physical Layer	La couche physique assure un transfert de bits sur le canal physique (support). À cet effet, elle définit les supports et les moyens d'y accéder : spécifications mécaniques (connecteur), spécifications électriques (niveau de tension), spécifications fonctionnelles des éléments de raccordement nécessaires à l'établissement, au maintien et à la libération de la ligne. Elle détermine aussi les moyens d'adaptation (ETCD ou DCE).
NIVEAU 2 Couche Liaison de données Data Link Layer	La couche liaison assure, sur la ligne, un service de transfert de blocs de données (trames) entre deux systèmes adjacents en assurant le contrôle, l'établissement, le maintien et la libération du lien logique entre les entités. Les protocoles de niveau 2 permettent, en outre, de détecter et de corriger les erreurs inhérentes aux supports physiques.
NIVEAU 3 Couche Réseau Network Layer	La couche réseau assure, lors d'un transfert à travers un système relais, l'acheminement des données (paquets) à travers les différents nœuds d'un sous-réseau (routage). Les protocoles de niveau 3 fournissent les moyens d'assurer l'acheminement de l'appel, le routage, le contrôle de congestion, l'adaptation de la taille des blocs de données aux capacités du sous-réseau physique utilisé. Elle peut offrir un service de facturation de la prestation fournie par le sous-réseau de transport.

COUCHES	FONCTIONS
NIVEAU 4 Couche Transport Transport Layer	La couche transport est la couche pivot du modèle OSI. Elle assure le contrôle du transfert de bout en bout des informations (messages) entre les deux systèmes d'extrémité (ETTD ou DTE). La couche transport est la dernière couche de contrôle des informations, elle doit assurer aux couches supérieures un transfert fiable quelle que soit la qualité du sous-réseau de transport utilisé.
NIVEAU 5 Couche Session Session Layer	La couche session gère l'échange de données (transaction) entre les applications distantes. La fonction essentielle de la couche session est la synchronisation des échanges et la définition de points de reprise.
NIVEAU 6 Couche Présentation Presentation Layer	Interface entre les couches qui assurent l'échange de données et celle qui les manipule, cette couche assure la mise en forme des données, les conversions de code nécessaires pour délivrer à la couche supérieure un message dans une syntaxe compréhensible par celle-ci. En outre, elle peut, éventuellement, réaliser des fonctions spéciales, comme la compression de données...
NIVEAU 7 Couche Application Application Layer	La couche application, dernière du modèle de référence, fournit au programme utilisateur, l'application proprement dite, un ensemble de fonctions (entités d'application) permettant le déroulement correct des programmes communicants (transferts de fichiers, courrier électronique...).

2.2 Principes de base d'une architecture en couche

Considérons le modèle simplifié à trois couches représenté figure 2.2.

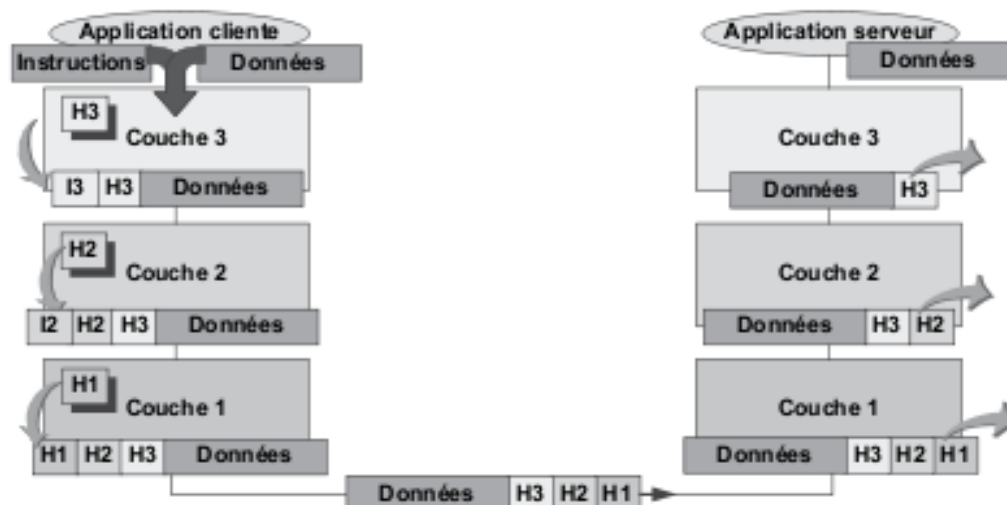


Figure 2.2 Principe général de fonctionnement d'un modèle en couches.

Pour communiquer, l'application cliente remet à la couche adjacente, ici la couche 3, des données à destination de l'application serveur, des instructions décrivant le service attendu et des informations nécessaires à l'acheminement des données vers l'application serveur. La couche 3 interprète les instructions reçues et confectionne une structure de données à destination de la couche 3 distante, dite couche homologue. Cette structure de données est constituée des données à transmettre auxquelles on ajoute un en-tête dit en-tête de niveau 3 (H3 pour *Header* de niveau 3) contenant un ensemble d'informations nécessaires à la couche 3 distante pour traiter les données. L'ensemble, en-tête et données, forme une unité de données de niveau N. Les règles d'échange entre couches de même niveau constituent un protocole de niveau N.

Puis, la couche 3 remet cette unité de données et des instructions (I3) à la couche inférieure qui procède de même... Enfin, les données sont émises sur le support physique. En réception la couche la plus basse extrait l'en-tête protocolaire (H1), l'interprète exécute les tâches demandées et remet les données à la couche supérieure qui procède de même jusqu'à remise des données à l'application distante. Le transport d'unité de données du niveau N dans une unité de données du niveau N-1 constitue ce que l'on désigne sous le terme d'encapsulation. Pour le niveau N les données encapsulées (données proprement dites et les divers en-têtes des niveaux inférieurs) constituent un ensemble d'octets sans signification. La figure 2.3 schématise pour les couches 1 à 4 le **principe de l'encapsulation**.

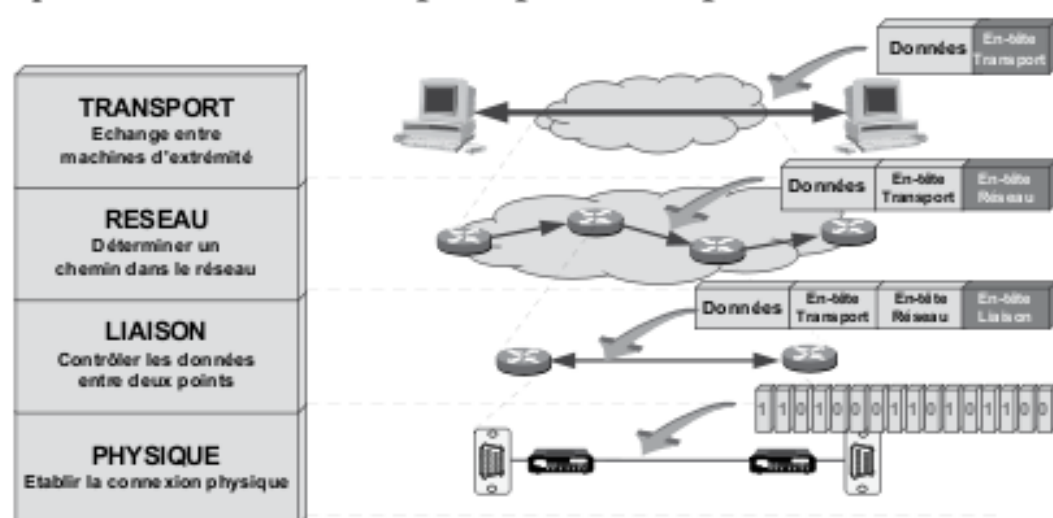


Figure 2.3 Principe de l'encapsulation.

3

L'architecture TCP/IP

3.1 Origine

L'architecture TCP/IP a été développée dans le milieu des années 1970 par la DARPA (*Defense Advanced Research Projects Agency* – États-Unis) pour les besoins de communication et d'interfonctionnement des applications entre les systèmes informatiques militaires (DoD, *Department of Defense*). Pour cela, il fallait définir un format d'échange des données commun à tous les systèmes tout en préservant l'existant, c'est-à-dire sans modifier les réseaux existants. En fait, TCP/IP masque aux applications les sous-réseaux réels de transport utilisés.

TCP/IP, du nom de ses deux protocoles principaux (TCP, Transmission Control Protocol et IP, Internet Protocol), est un ensemble de protocoles permettant de résoudre les problèmes d'interconnexion en milieu hétérogène. À cet effet, TCP/IP décrit un réseau logique (réseau IP) au-dessus du ou des réseaux physiques réels qui réalisent le transport effectif des données et auxquels sont effectivement connectés les ordinateurs (figure 3.1).

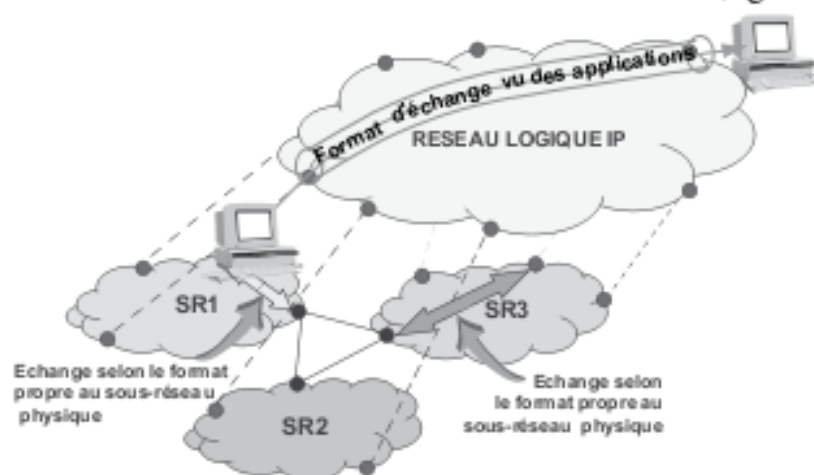


Figure 3.1 Le réseau logique IP et sous-réseaux physiques réels (SRx).

Dans cette approche, les échanges entre applications sont réalisés selon le format défini par TCP/IP, alors que l'échange des données dans les sous-réseaux physiques réels se réalise selon le format propre à chaque sous-réseau.

3.2 Principe architectural

Précédant le modèle OSI, TCP en diffère fortement, non seulement par le nombre de couches, mais aussi par l'approche. Le modèle OSI spécifie des services (approche formaliste), TCP/IP des protocoles (approche pragmatique). Développé au-dessus d'un environnement existant, TCP/IP ne décrit, à l'origine, ni de couche physique ni de couche liaison de données. Les applications s'appuient directement sur le service de transport. Aussi l'architecture TCP/IP de base ne comprenait que deux couches : la couche transport (TCP) et la couche inter-réseau (IP). La figure 3.2 compare les deux architectures.

Il n'y a pas de couche application au sens OSI du terme, c'est-à-dire de couche présentant des « API » (*Application Programming Interface*) aux applications qui rendent transparents à ces dernières le ou les sous-réseaux réels de transport utilisés. Cependant, un mécanisme particulier, les sockets, assure une communication d'application à application en masquant les éléments réseaux.

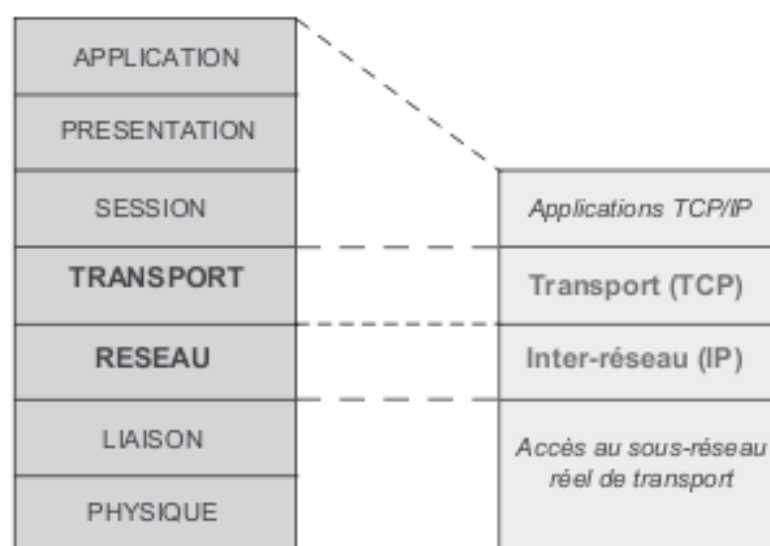


Figure 3.2 Le modèle OSI et l'architecture TCP/IP.

3.3 Description générale de l'environnement TCP/IP

L'architecture TCP/IP comprend de nombreux programmes applicatifs, utilitaires et protocoles complémentaires (figure 3.3). À l'origine TCP/IP ne spécifiait aucun protocole de liaison, il s'appuyait sur les réseaux existants. L'utilisation massive de TCP/IP a fait apparaître le besoin de liaisons tout IP et donc la nécessité de disposer de protocoles de liaison spécifiques (SLIP, PPP). De même, TCP/IP a été adapté aux protocoles dits « haut débit » comme le Frame Relay et l'ATM (*Asynchronous Transfer Mode*), ce dernier constituant encore aujourd'hui le cœur de la plupart des réseaux privés et d'opérateurs.

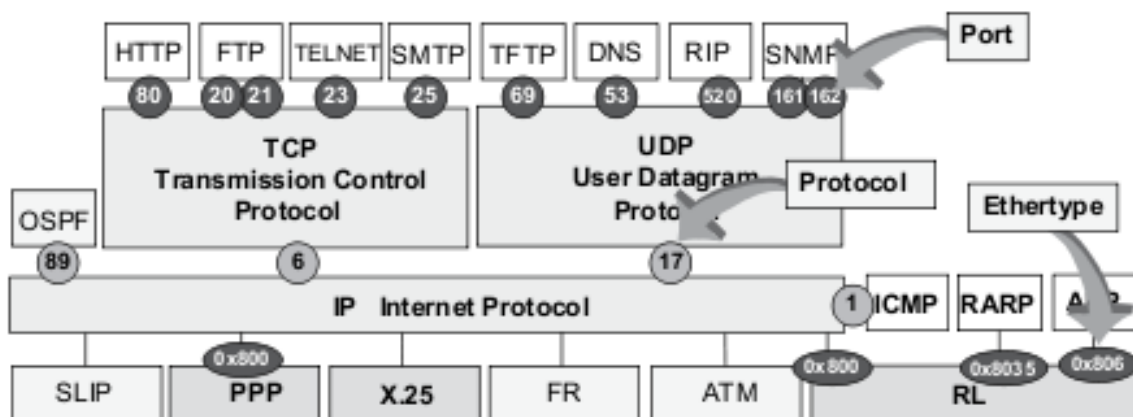


Figure 3.3 Les protocoles et les applications de TCP/IP.

Les principaux protocoles et applications de l'environnement TCP/IP qui seront décrits dans les chapitres suivants sont :

- ▶ **ARP**, *Address Resolution Protocol*, met en correspondance une adresse logique IP avec une adresse physique **MAC** (*Medium Access Control*, adresse de l'interface dans les réseaux locaux) ;
- ▶ **DNS**, *Domain Name System*, est un système de bases de données réparties assurant la correspondance entre un nom symbolique et une adresse internet (adresse IP) ;

- ▶ **FTP**, *File Transfer Protocol*, est un système de manipulation de fichiers à distance (transfert, suppression, création...) ;
- ▶ **HTTP**, *HyperText Transport Protocol*, assure le transfert de fichiers hypertextes entre un serveur web et un client web ;
- ▶ **ICMP**, *Internet Control and error Message Protocol*, assure un dialogue IP/IP et permet notamment : la signalisation de la congestion, la synchronisation des horloges et l'estimation des temps de transit... Il est utilisé par l'utilitaire **Ping** qui permet de tester la présence d'une station sur le réseau.
- ▶ **OSPF**, *Open Shortest Path First*, est un protocole de routage du type état des liens, il a succédé, dans le réseau Internet, au protocole RIP ;
- ▶ **PPP**, *Point to Point Protocol*, protocole d'encapsulation des datagrammes IP, il assure la délimitation des trames, identifie le protocole transporté et assure la détection d'erreurs.
- ▶ **RARP**, *Reverse Address Resolution Protocol*, permet l'attribution d'une adresse IP à une station ;
- ▶ **RIP**, *Routing Information Protocol*, est le premier protocole de routage (vecteur distance) utilisé dans Internet ;
- ▶ **SLIP**, *Serial Line Interface Protocol*, protocole d'encapsulation des paquets IP, ce protocole n'assure que la délimitation des trames ;
- ▶ **SMTP**, *Simple Mail Transfer Protocol*, offre un service de courrier électronique ;
- ▶ **SNMP**, *Simple Network Management Protocol*, est devenu le standard des protocoles d'administration de réseau ;
- ▶ **TELNET**, *TELEtypewriter NETwork protocol* (ARPA) ou *TERminal NETwork protocol*, système de terminal virtuel, permet l'ouverture de sessions avec des applications distantes ;
- ▶ **TFTP**, *Trivial FTP*, est une version allégée du protocole FTP.

Afin de distinguer le protocole ou l'application auquel doivent être remises les données, il a été défini, à l'instar du modèle OSI avec la notion de SAP (*Service Access Point*), un adressage de couche, cet identifiant est

transporté dans l'en-tête protocolaire ; ainsi, par exemple, l'EtherType des trames « Ethernet¹ » identifie le protocole du niveau réseau. L'identifiant de Protocole dans le datagramme IP désigne le protocole de transport utilisé et la notion de Port dans le segment TCP détermine l'instance locale de l'application. La figure 3.3 illustre ce principe et donne quelques exemples d'identifiants normalisés.

3.4 Conclusion

L'intégration de TCP/IP à UNIX BSD 4, par l'université de Berkeley, en fit le standard de la communauté UNIX (1980). En 1983, TCP/IP a remplacé le protocole NCP (*Network Control Program*) dans ARPANET, ancêtre de l'Internet. Aujourd'hui, TCP/IP est le protocole adopté dans tous les réseaux, du LAN au WAN. A l'origine conçu pour des applications en mode texte, TCP/IP a dû s'adapter pour répondre aux exigences des nouvelles applications (figure 3.4) données/voix/image notamment par la mise en œuvre de mécanisme de qualité de service (QoS).

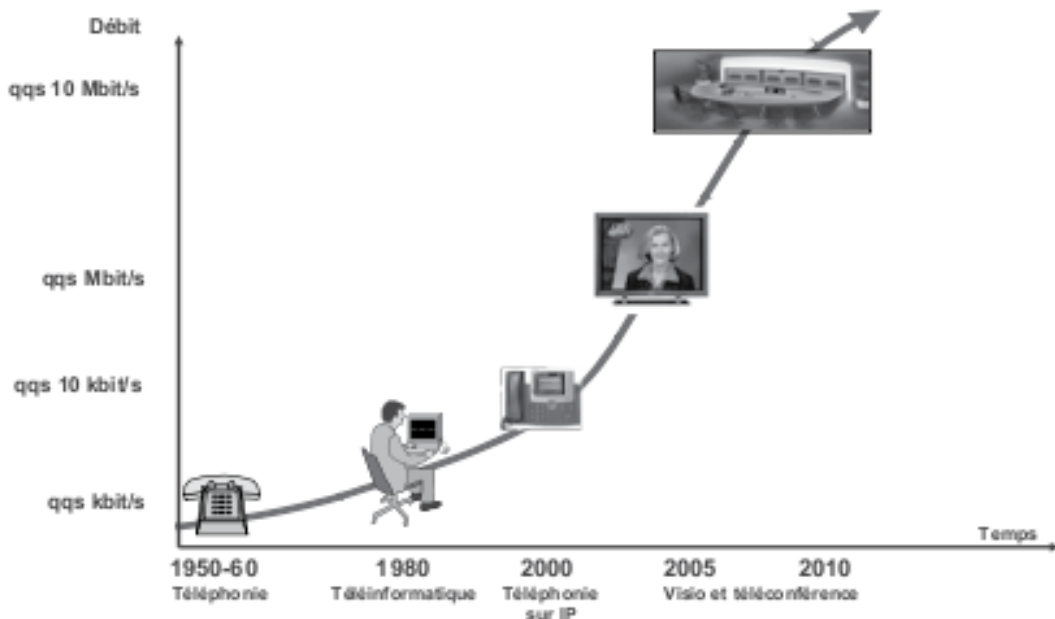


Figure 3.4 L'évolution des services offerts par les réseaux.

¹ Ethernet est le nom de marque déposé par Xerox d'un type de réseau local ; ce terme, passé dans le langage courant, désigne par abus de langage tous les réseaux locaux utilisant le protocole d'accès CSMA/CD.

2

Éléments physiques de la liaison de données

4

Les supports et leur limitation

L'infrastructure d'un réseau, la qualité de service offerte, les solutions logicielles à mettre en œuvre dépendent largement des supports de transmission utilisés. Les supports de transmission exploitent les propriétés de conductibilité des métaux (paires torsadées, câble coaxial), celles des ondes électromagnétiques (faisceaux hertziens, guides d'ondes, satellites) ou encore celles du spectre visible de la lumière (fibre optique). Généralement on classe les supports en deux catégories :

- ▶ les supports guidés (supports cuivre et supports optiques) ;
- ▶ les supports libres (faisceaux hertziens et liaisons satellites).

4.1 Les supports cuivre

Un support de transmission quel qu'il soit est généralement qualifié par sa bande passante et les supports cuivre par leur impédance caractéristique et leur résistance ohmique, source de perte par effet Joule.

4.1.1 La bande passante et le système de transmission

L'impulsion électrique (bit) est un phénomène discontinu qui ne peut être modélisé. Le mathématicien et physicien Fourier a montré que tout signal périodique non sinusoïdal (par assimilation une suite de bits 01010101) peut être considéré comme la somme d'une composante continue (A_0) et d'une infinité de signaux sinusoïdaux d'amplitude, de fréquence et de phase convenablement choisies, comme l'illustre la figure 4.1.



$$u(t) = 4 U / \pi (\sin \omega t + 1/3 \sin 3\omega t + 1/5 \sin 5\omega t + \dots)$$

Figure 4.1 La décomposition d'un signal carré symétrique.

Un signal périodique quelconque est donc composé d'une infinité de signaux sinusoïdaux. L'espace de fréquences occupé par les composantes se nomme **largeur de bande**. En théorie, la largeur de bande d'un signal non sinusoïdal est infinie. Un système ne peut, en aucun cas, transmettre toutes ces composantes, son espace de réponse en fréquences est limité. La **bande passante** d'un système exprime l'espace de fréquences qu'un système peut correctement transmettre sans que les signaux ne subissent un affaiblissement de plus de la moitié de leur puissance d'origine, c'est la bande passante¹ définie généralement -3 dB (figure 4.2).

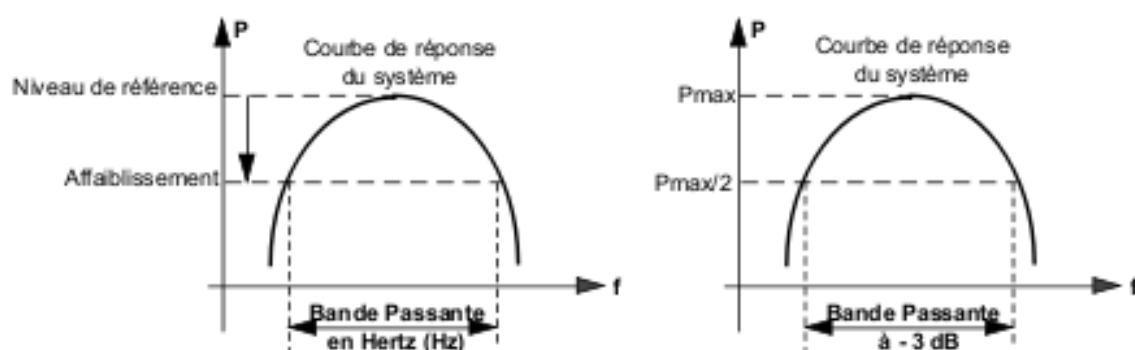


Figure 4.2 La bande passante à - 3 dB.

À l'extrémité de la ligne, le récepteur doit identifier et décoder le signal. Cette fonction ne peut valablement être réalisée que si le signal n'a pas été exagérément altéré durant la transmission. Ces modifications dépendent

¹ L'affaiblissement, exprimé en décibel (dB), est donné par la relation :

$$A = 10 \log_{10} P_1 / P_0$$

où P_1 est la puissance du signal en sortie,
 P_0 est la puissance du signal de référence

d'une part de la nature du signal (largeur de bande ou encore spectre du signal) et, d'autre part, de la réponse en fréquence du système (bande passante). Les systèmes de transmission (lignes, amplificateurs...) ne transmettant pas toutes les composantes du signal de façon identique déforment les signaux, c'est la **distorsion** (figure 4.3). Ainsi, dans un système réel, les signaux sont transmis avec une distorsion faible jusqu'à une certaine fréquence appelée **fréquence de coupure**. Au-delà de cette fréquence, toutes les composantes sont fortement atténuées.

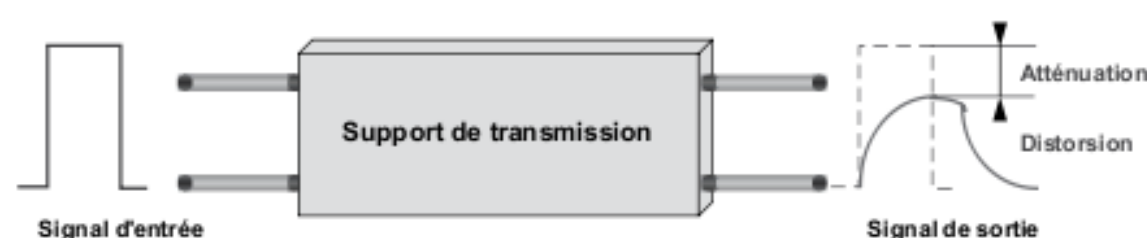


Figure 4.3 La déformation du signal par le support de transmission.

La **largeur de bande** d'un signal correspond à la bande passante minimale que le système doit posséder pour interpréter valablement le signal et restituer correctement l'information. Ainsi, la bande passante qualifie le système, et la largeur de bande qualifie le signal. Notons que le terme de bande passante est utilisé non seulement pour désigner un espace fréquentiel (bande passante ou BP en Hz), mais aussi pour qualifier le débit binaire d'un système (bande passante exprimée en bit/s).

4.1.2 L'impédance caractéristique et l'adaptation d'impédance

Une ligne de transmission présente au courant électrique un effet résistif (R) responsable de l'atténuation du signal, des effets réactifs qui se décomposent en effet selfique (L) et en effet capacitif (C), et enfin la conductance (G) qui exprime la perte par effet résistif entre les deux conducteurs (généralement négligeable). La notion d'impédance en courant alternatif recouvre une notion similaire à celle de résistance en courant continu mais une « résistance » qui varierait en fonction de la fréquence du courant qui la traverse. L'impédance, comme la résistance, s'exprime en ohm (Ω).

Le rapport u/i ¹ pour une ligne supposée de longueur infinie s'appelle impédance caractéristique notée Z_c .

On montre qu'une ligne de longueur finie refermée sur un récepteur, dont l'impédance Z_r serait telle que $Z_r = Z_c$, se comporte comme une ligne de longueur infinie. Le transfert de puissance est alors maximal entre le générateur et le récepteur, la ligne est dite adaptée (adaptation d'impédance). Lorsque deux systèmes d'impédance différente sont mis en relation, le transfert de puissance n'est pas optimal, la puissance non absorbée par le système distant est réfléchiée vers la source, ce phénomène s'appelle l'écho. En transmission numérique, l'écho a pour conséquence de générer des « bits fantômes », introduisant ainsi des erreurs de transmission. La figure 4.4 illustre un système complètement désadapté. À chaque point de raccordement, une partie de l'énergie est réfléchiée. La source reçoit deux fois le signal d'écho. Le premier dû à la rupture d'impédance locale (écho local) est peu gênant. Le second dû à la rupture d'impédance distante (écho distant) est plus perturbant, des dispositifs spécifiques (annuleurs d'écho) ont en charge de supprimer les effets de cet écho.

Pour éviter ces réflexions parasites, il est nécessaire, tout au long de la ligne et à chaque raccordement d'un nouvel élément à la liaison, de réaliser la continuité de l'impédance : c'est l'adaptation d'impédance.

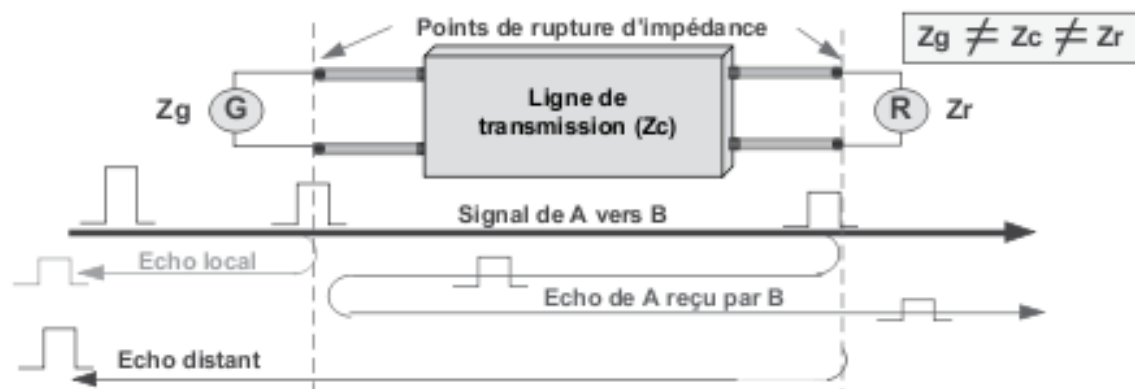


Figure 4.4 La notion d'écho.

1 Variation de la tension (u) sur celle du courant (i).

4.2 Les supports guidés

4.2.1 La paire torsadée

La paire torsadée ou symétrique est constituée de deux conducteurs identiques torsadés. Généralement, plusieurs paires sont regroupées sous une enveloppe protectrice appelée gaine pour former un câble. Les câbles contiennent une paire (desserte téléphonique), quatre paires (réseaux locaux), ou plusieurs dizaines de paires (câble téléphonique). La figure 4.5 illustre les principaux types de paires torsadées utilisés dans les réseaux locaux.

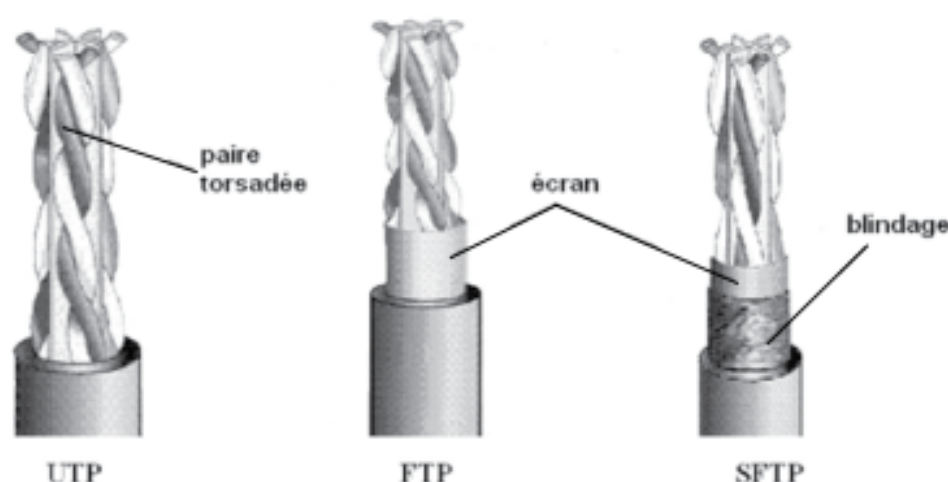


Figure 4.5 Les différentes paires torsadées (Catégorie 6).

La paire torsadée « ordinaire » (paire symétrique, UTP *Unshielded Twisted Pairs*) est sensible à l'environnement électromagnétique (parasites industriels, proximité de câbles à courant fort...). L'utilisation de tels câbles est soumise à des contraintes d'installation. Pour ces raisons, la paire symétrique est généralement utilisée sans référence à la terre (transmission différentielle), ce qui améliore sa résistance aux parasites.

L'immunité aux parasites peut être améliorée en protégeant le faisceau par un écran (câble écranté). L'écran est constitué d'un ruban d'aluminium qui entoure les paires et les protège des perturbations électromagnétiques. Un conducteur de cuivre nu étamé (drain) permet la mise à la terre de l'écran (paires écrantées, FTP *Foiled Twisted Pairs*). Une meilleure protection peut encore être obtenue en réalisant, autour de chacune des paires, un véritable blindage (paires blindées, STP *Shielded Twisted Pairs*).

Tableau 4.1 Les catégories de paires torsadées.

Catégorie	Classe	Impédance	Fréquence max.	Applications
3	C	100-120 Ω	16 MHz	Token Ring 4 Mbit/s 10 Base T Fast Ethernet 100 VG AnyLAN 100 Base T4
4	D	100 Ω	20 MHz	Token Ring 16 Mbit/s
5	D	100 Ω	100 MHz	Câble UTP et FTP 100 Base Tx ATM 155 Mbit/s 1000 Base T (Cat 5E)
6	E	100 Ω	250 MHz	Câble UTP, FTP et SFTP 1 000 Base Tx
6a	E	100 Ω	5 800 MHz	Câble UTP, FTP et SFTP 1 000 Base Tx 10 G Base T
7	F	100 Ω	600 MHz	Câble SFTP

La paire symétrique est actuellement le conducteur le plus utilisé : dessert locale des raccordements téléphoniques, liaisons d'accès aux réseaux de données et surtout les réseaux locaux où les faibles distances autorisent l'utilisation de débits élevés : 100 Mbit/s sur 100 m, voire 10 Gbit/s sur 100 m avec de la paire UTP. Le tableau 4.1 indique pour chaque catégorie de paires torsadées sa principale utilisation.

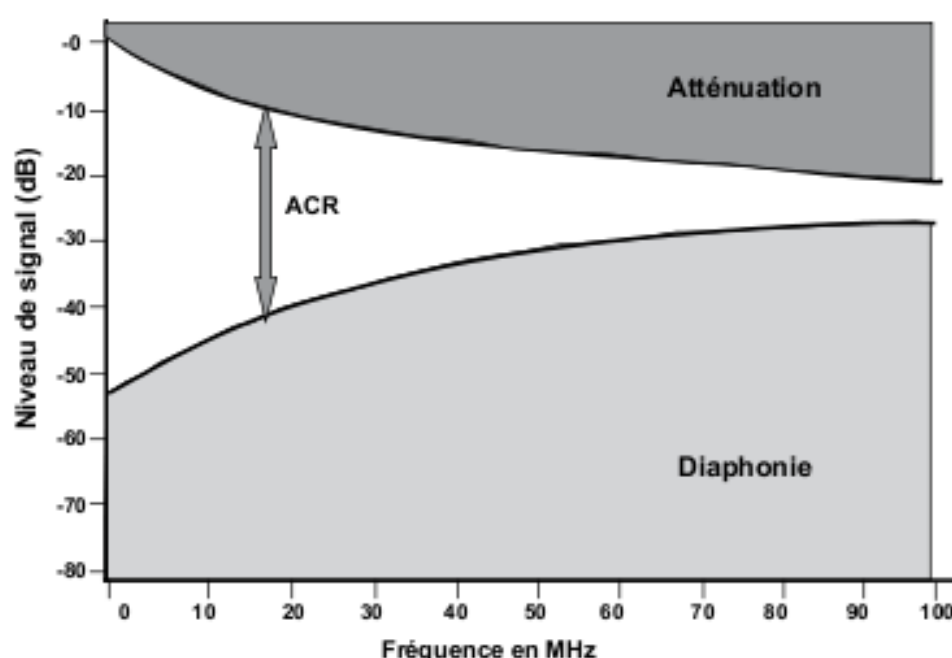


Figure 4.6 Relation entre la diaphonie et le signal (ACR).

La catégorie distingue les équipements, la notion de classe de câblage qualifie un câblage de bout en bout. Les principales caractéristiques des câbles à paires torsadées sont représentées figures 4.6 et 4.7 :

- ▶ L'impédance caractéristique (Z_c), aujourd'hui seuls les câbles 100 Ω sont utilisés dans les réseaux locaux ;
- ▶ L'affaiblissement ou atténuation mesuré en dB ;
- ▶ Les interférences internes au câble dues à un seul émetteur :
 - Le **Return Loss** (RL), ou affaiblissement de réflexion, mesure le rapport entre l'énergie émise par la source et celle reçue par celle-ci due aux réflexions, provoquées par les ruptures d'impédance de la ligne, de son propre signal ;
 - Le **NEXT** (*Near end crosstalk*) ou diaphonie locale (paradiaphonie) entre deux paires du même côté, plus la valeur est importante, meilleur est le câble. Cette valeur dépend du pas de torsade et de la régularité de celui-ci ;
 - L'**ACR** (*Attenuation crosstalk ratio*) qui représente la différence entre l'atténuation et la diaphonie, c'est un rapport entre la puissance du signal et celle des interférences ;

- Le **FEXT** (*Far end crosstalk*) mesure le rapport entre la télédiaphonie et l'affaiblissement. Plus la valeur est grande, meilleur est le canal ;
- Les interférences internes induites sur une paire par l'émission de toutes les autres paires :
 - Le **PSNEXT** (*Power sum near-end crosstalk*), ou paradiaphonie cumulée, mesure la perturbation de toutes les paires sur une seule, cette valeur ne résulte pas d'une mesure mais d'un calcul ;
 - Le **PSFEXT** (*Power-sum far-end crosstalk*) diaphonie totale distante mesure l'effet cumulé d'une diaphonie distante sur une paire en provenance de toutes les autres paires du câble ;
 - L'**AXT** (*Alein crosstalk*), ou diaphonie exogène, mesure l'ensemble des interférences résultant de tous les câbles du toron sur une paire. L'AXT ne concerne que les câbles UTP ;
- Le **delays skew** ou différence de propagation mesure la différence de temps de propagation entre la paire la plus rapide et la moins rapide, il s'exprime en nanoseconde ;
- **NVP** (*Nominal Velocity of Propagation*) ou coefficient de vélocité, c'est le rapport entre la vitesse de propagation dans le câble et la vitesse de la lumière, il est exprimé en % de la célérité.

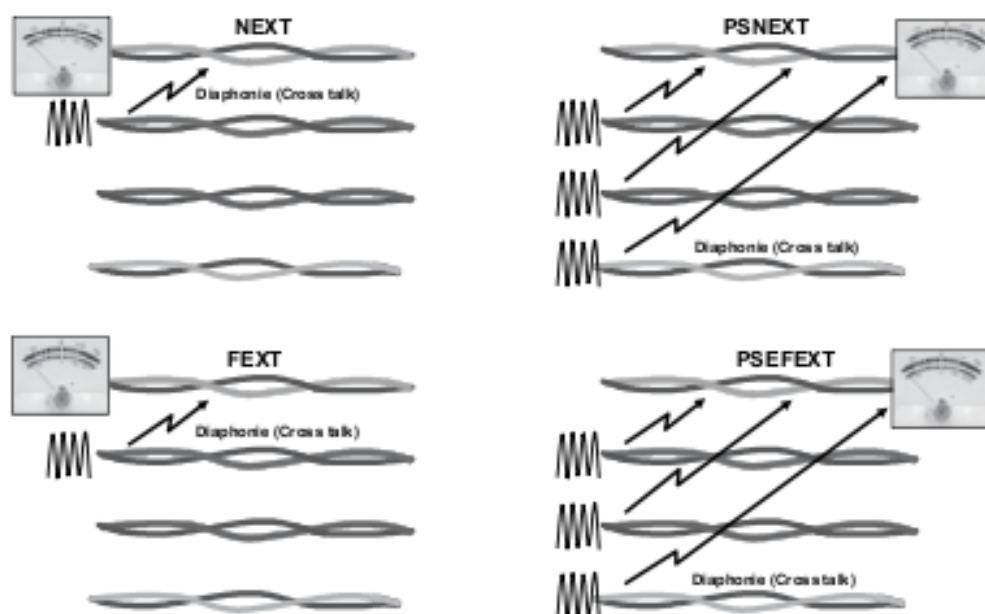


Figure 4.7 Mesures de la diaphonie.

4.2.2 La paire torsadée dans les réseaux (câblage)

■ Le précâblage d'immeuble

Le développement intensif des postes de travail en réseau local a révélé des problèmes liés au câblage. Les réseaux locaux ont tous, aujourd'hui, une topologie physique en étoile, d'où l'idée de réaliser, dans les immeubles de bureaux, un précâblage (figure 4.8). Un système de précâblage comporte une distribution horizontale et si le réseau couvre plusieurs étages un câblage vertical souvent réalisé en fibre optique, la distribution horizontale doit :

- ▶ assurer que tout poste de travail n'est qu'à quelques mètres d'une prise informatique ou téléphonique ;
- ▶ être indépendant du type de réseau et de la topologie logique du réseau choisi.

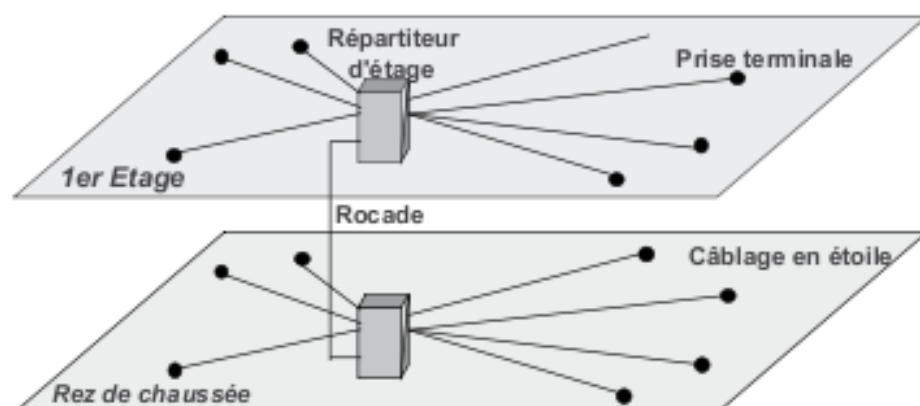


Figure 4.8 La structure générale d'un précâblage d'immeuble.

Différents systèmes de câblage ont été définis par les principaux constructeurs, ce sont principalement l'ICS (*IBM Cabling System*), le BCS (*Bull Cabling System*), l'Open Link de DEC et le PDS *Systimax* d'origine AT&T. Ces systèmes ont en commun l'utilisation de la paire torsadée et une topologie physique en étoile. Le cœur du câblage est constitué de panneaux, dits **panneaux de brassage**, qui permettent, à l'aide de jarretières, de réaliser la connexion des postes de travail selon la topologie requise par le réseau. Ces systèmes diffèrent essentiellement par le type de câble utilisé (UTP, FTP...).

Les câbles de catégorie 5 et 5e sont aujourd'hui les plus installés. Les dernières installations concernent essentiellement la catégorie 6 et 6a. La figure 4.9 représente les différents constituants participant à la réalisation physique d'un câblage en paires torsadées. Un local technique (local de brassage) abrite une armoire technique ou armoire de brassage. Celle-ci accueille les éléments actifs (hub, MAU ou *Medium Access Unit*), parfois les serveurs (salle réseau). Le précâblage consiste à « tirer » des câbles entre le panneau de brassage et les prises murales. Pour raccorder une station, il suffit de connecter celle-ci à la prise murale par un cordon dit de raccordement. À l'autre extrémité, au panneau de brassage, on tire une jarretière (cordon de raccordement ou de brassage) entre la prise RJ45 du panneau de brassage correspondant à l'extrémité de la prise murale à activer et un port de l'élément actif (hub, commutateur...). Chaque point de raccordement pouvant constituer une désadaptation d'impédance et engendrer des ondes stationnaires, il est impératif de veiller à la qualité de tous ces éléments et de vérifier qu'ils sont qualifiés pour le type de réseau mis en œuvre. L'erreur la plus fréquemment constatée étant l'utilisation de jarretières UTP sur un câblage FTP ou l'inverse.

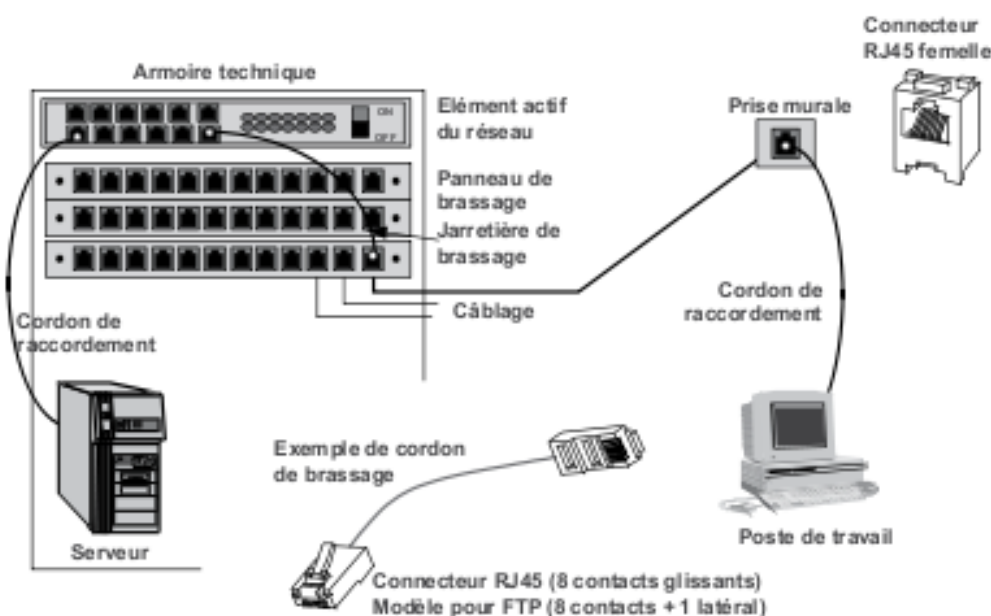


Figure 4.9 Les constituants physiques d'un réseau local.

Les seules contraintes d'installation de la paire torsadée concernent la distance minimale à respecter entre le câble et les sources de rayonnement

électromagnétique (distribution courant fort, tubes fluorescents, machinerie électrique...). Par exemple, le cheminement parallèle de câbles courant faible et courant fort doit respecter un écartement minimal de 5 cm si le cheminement est inférieur à 10 m, 15 cm jusqu'à 30 m et 30 cm au-delà. Le rayon de courbure des câbles (8 fois de diamètre du câble pour le FTP, 4 fois pour l'UTP) et la longueur maximale de « détorsadage » des paires pour réaliser la connectique (1,5 cm) sont des points à examiner lors de la recette d'un câblage.

■ La téléalimentation (IEEE 802.3af et at)

Au-delà de la téléphonie sur IP, la recommandation IEEE 802.3af et son évolution 802.3at (**PoE**, *Power over Ethernet*) permettent la télé-alimentation des points d'accès *Wireless* (réseau sans fil), des caméras auto-alimentées...

La technologie PoE fournit aux équipements terminaux une puissance électrique maximale de 45W sous une tension de 48 V continu (CC). Deux modes d'injection sont décrits, l'énergie peut être fournie directement par l'élément actif (commutateur), ce mode est dit *End-Span* ou par un équipement supplémentaire intermédiaire (mode *Mid-Span*). Cette dernière technique dite par panneaux de brassage utilise un équipement spécifique d'aspect similaire à un commutateur, elle nécessite un double brassage (figure 4.10) et de disposer de place dans les baies de brassage.

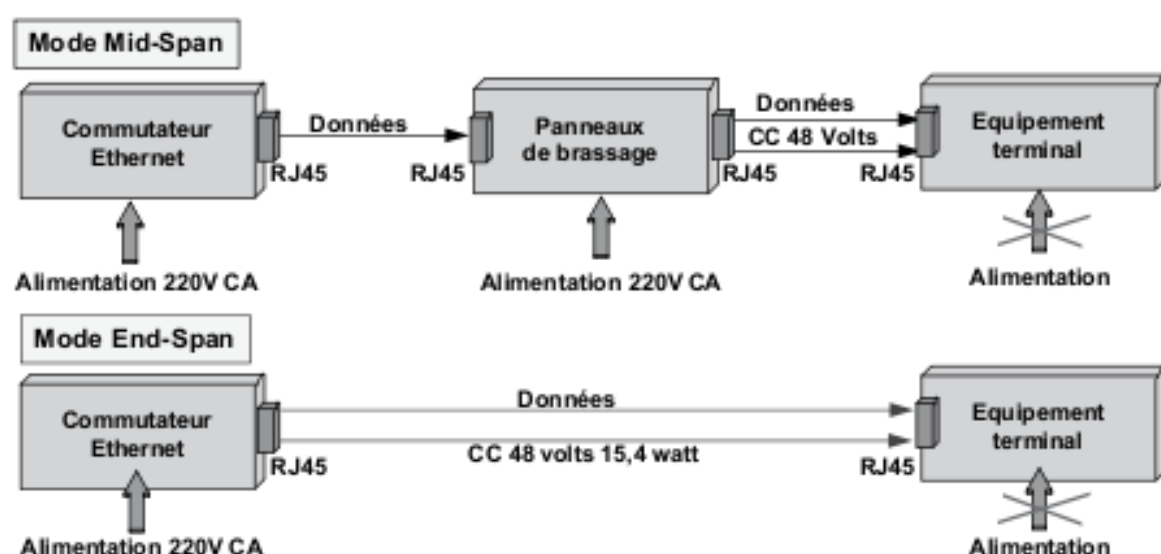


Figure 4.10 Les modes de télé-alimentation.

Afin de ne pas endommager l'équipement terminal, le système d'injection (PSE, *Power Sourcing Equipment*) doit s'assurer que l'équipement terminal est compatible 802.3af. Les PSE implémentent un système de détection. Ce dernier vérifie que l'équipement terminal peut admettre la télé-alimentation et détermine quelles sont les paires utilisées pour fournir l'énergie électrique¹.

La télé-alimentation peut être fournie par les paires non utilisées pour la transmission des données (*Spare Pairs*, paires 1 et 4 disponibles) ou sur les paires de données (*Signal Pair*, paires 2 et 3). Ces deux modes sont représentées en figure 4.11.

Le tableau 4.2 indique pour la norme IEEE 802.3at la puissance injectée (PSE, *Power Supply Equipment*) et celle minimale disponible pour l'équipement terminal (PD, *Power Device*) :

Tableau 4.2 Puissance disponible IEEE 802.3at.

Type de PSE	Classification	Puissance injectée	Puissance disponible
Type 1	0	15,4	13
	1	4,0	3,84
	2	7,0	6,49
	3	15,4	13
Type 2	4	30	25,5

Dans la pratique, beaucoup de commutateurs n'offrent pas une puissance disponible qui permette de supporter la classe 4 sur tous les ports. Certains constructeurs mettent en œuvre un logiciel de répartition de la puissance, ce qui contraint à vérifier l'adéquation entre le bilan électrique nécessaire et celui disponible sur l'équipement.

¹ Le protocole LLDP (Link Layer Discovery Protocol, IEEE 802.A) reconnaît le type de terminal et l'éventuelle auto-alimentation.

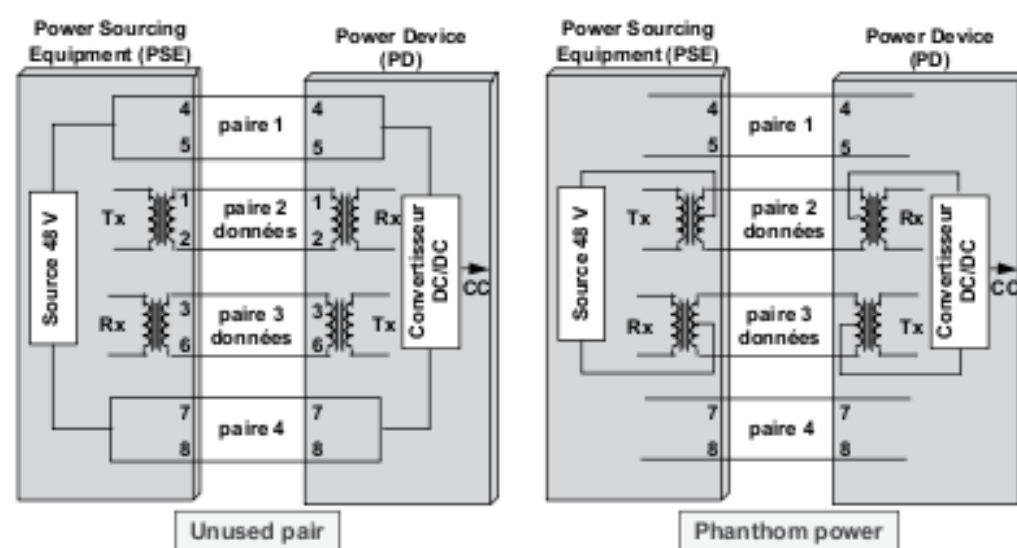


Figure 4.11 Les deux modes de télé-alimentation.

■ Le câblage et son environnement thermique

Pour l'application de la norme IEEE 802.3at, la recommandation TR42.7 fixe un courant maximal de 720 mA par paire et une température ambiante maximale de 45 °C, température au-delà de laquelle les performances du câble sont fortement altérées et son vieillissement accéléré. Lors de l'étude d'un nouveau câblage, la notion de température maximale devient une contrainte forte dans la définition du nombre de câbles maximal par toron et du cheminement de ceux-ci (proximité des éléments de chauffage). Le tableau 4.3 indique, par type de câble, l'élévation de température à prendre en compte. Si on admet une température ambiante de 20 °C, l'application de cette recommandation limite chaque toron à une cinquantaine de câbles.

Tableau 4.3 IEEE 802.3at et l'élévation de température.

Jauge	Diamètre câble	Résistance pour 100m	Puissance dégagée 802.3at	Élévation de température
AWG 26	0,405 mm	13,20 Ω	1,71 W	0,56 °C
AWG 24	0,511 mm	8,29 Ω	1,07 W	0,22 °C
AWG 23	0,573 mm	6,59 Ω	0,85 W	0,14 °C
AWG 22	0,674 mm	5,22 Ω	0,68 W	0,09 °C

■ Mise à la terre des câblages de type FTP

La différence visible entre un câble UTP et FTP est évidemment l'écran de protection, cette différence physique de constitution du câble induit des contraintes d'installation et particulièrement sur la distribution des terres. Le câble UTP n'implique aucune contrainte spécifique sur les terres, ce n'est pas le cas du câblage FTP (figure 4.12). La norme Ethernet IEEE 802.3 spécifie que l'isolation galvanique entre les systèmes doit être au minimum de 1 500 Volts et précise que les masses des équipements ne doivent pas être reliées entre elles par le circuit de télécommunication, ce que seul le câblage UTP respecte.

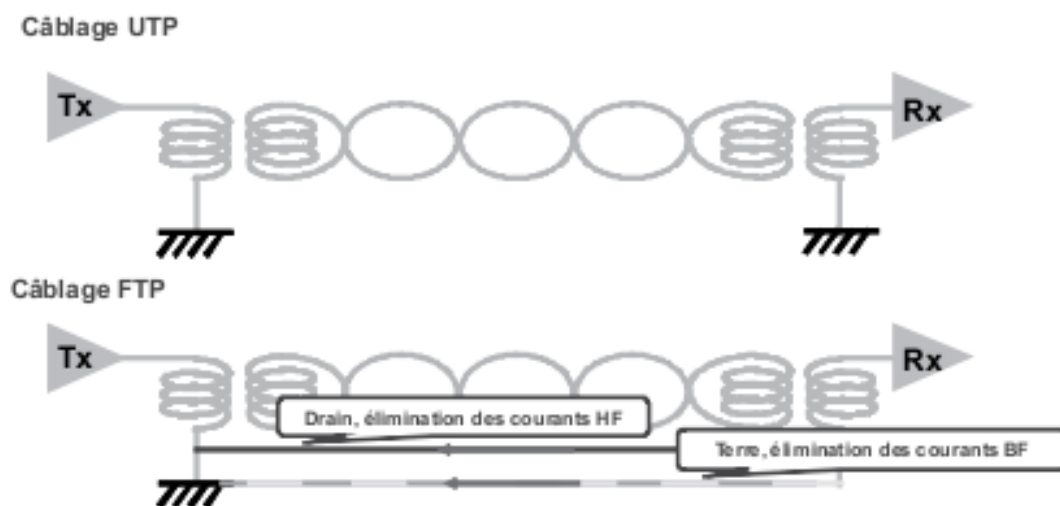


Figure 4.12 Mise à la terre des câblages UTP/FTP.

L'obligation de fusionner les terres électrique et informatique (plus qu'une seule distribution de terre, Norme NF EN 50174) a introduit des conditions drastiques sur les spécifications d'un câblage FTP¹ telles que celles-ci sont difficilement voire jamais respectées et d'ailleurs ne peuvent l'être avec les équipements télé-alimentés qui ne présentent aucune référence à la terre (figure 4.13).

¹ La norme ISO 11801 du 23 octobre 2002 reprise pour harmonisation dans la norme européenne EN 50174 (NF 50174) spécifie qu'un câble FTP doit avoir ses deux extrémités reliées à la terre (drain au panneau de brassage et *via* l'équipement terminal du côté utilisateur). À cette contrainte, la norme ISO 11801 en ajoute une seconde : la différence de potentiel entre les deux points de masse doit rester inférieure à un volt RMS.

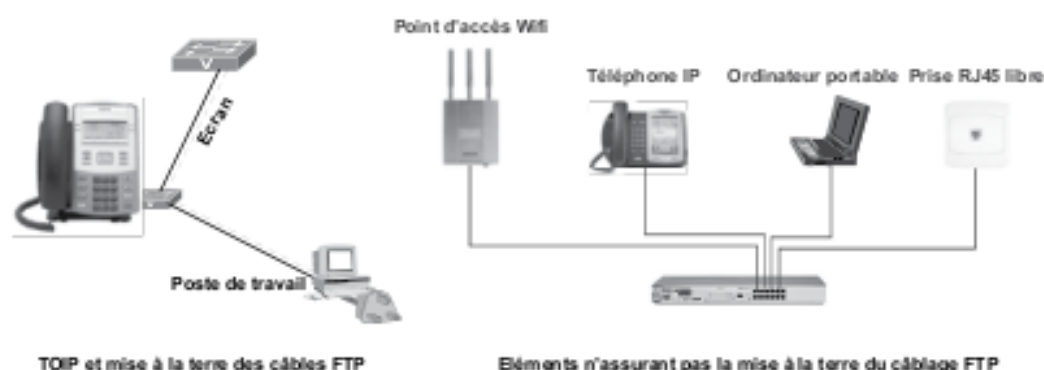


Figure 4.13 Problématique de la mise à la terre.

Aujourd'hui, 83 % des installations de câblage respectent cette normalisation issue de l'IEEE, mais, de par l'utilisation d'écrans ou de blindages reliant les masses des équipements actifs de réseaux 17 % des câblages, soit environ 80 % des câblages français sont en contradiction avec la norme Ethernet (données mondiales).

4.2.3 Le câble coaxial

Une paire coaxiale ou câble coaxial (figure 4.14) est constituée de deux conducteurs concentriques maintenus à distance constante par un diélectrique. Le conducteur extérieur, relié à la terre, est constitué d'une tresse métallique en cuivre recuit appelée **blindage**. L'ensemble est protégé par une gaine isolante.

Le câble coaxial, couramment appelé câble CATV (câble télévision), possède des caractéristiques électriques supérieures à celles de la paire torsadée. Il autorise des débits plus élevés et est peu sensible aux perturbations électromagnétiques extérieures. Le taux d'erreur sur un tel câble est d'environ 10^{-9} .

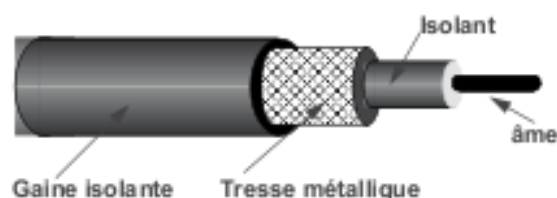


Figure 4.14 Le câble coaxial.

Le CATV présente une bonne immunité aux parasites, mais il est cher et exigeant en contraintes d'installation (rayon de courbure...), il n'est plus

utilisé que dans des environnements perturbés ou dans les systèmes sécurisés (rayonnement). Dans les réseaux locaux, il est, aujourd'hui, remplacé par la paire torsadée et dans les liaisons longues distances par la fibre optique.

4.2.4 La fibre optique

■ Principe

Un faisceau de lumière (figure 4.15), au passage d'un milieu 1 vers un milieu 2 (dioptre), est réfléchi (retour au milieu d'origine) et est réfracté avec une déviation (passage dans le milieu 2).

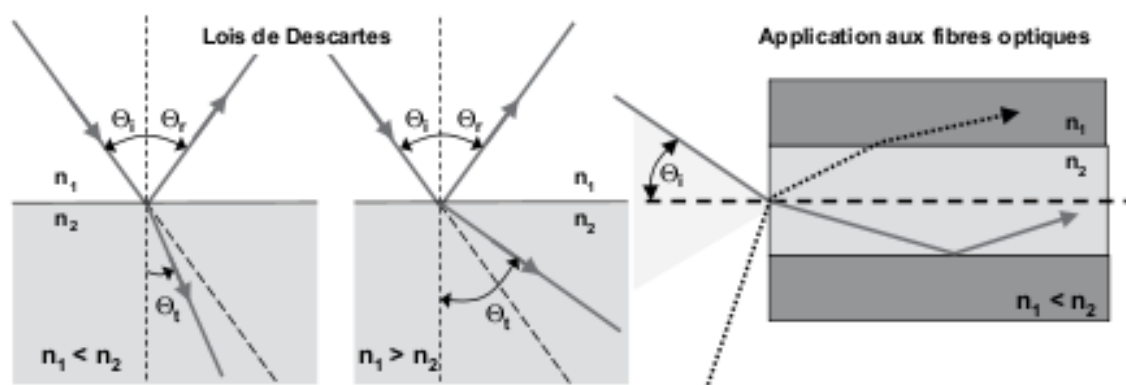


Figure 4.15 Lois de Descartes, application aux fibres optiques.

L'indice de réfraction (n_1, n_2) mesure le rapport entre la vitesse de propagation du rayon lumineux dans le vide et celle dans le milieu considéré, soit :

$$n = c / v$$

où n est l'indice de réfraction absolu du milieu considéré,
 c la vitesse de la lumière dans le vide ($3 \cdot 10^8$ m/s),
 v vitesse de propagation de la lumière dans le milieu considéré.

L'indice de réfraction du vide est de 1, celui de l'air 1,003, du verre ordinaire d'environ 1,5 et de l'eau 1,33.

Lorsque l'angle d'incidence augmente (θ_i), l'énergie réfractée diminue et l'énergie réfléchie augmente. Si on augmente encore l'angle, la réfraction

devient nulle ($\theta_2 = \pi/2$, condition limite de la réfraction) toute l'énergie est réfléchi (réflexion totale). Cette propriété est utilisée pour réaliser des guides de lumière : la fibre optique (figure 4.15).

Une fibre optique (figure 4.16) est composée d'un « fil » de silice étiré de telle manière que l'on distingue deux structures d'indice de réflexion différent, l'une appelée **cœur**, l'autre qui « l'entoure » appelée gaine optique ou **manteau**, l'ensemble étant protégé par une enveloppe dite de protection. La lumière s'y propage par réflexions successives à la frontière de la rupture d'indice de réflexion, la réflexion totale est assurée par des valeurs d'indices proches, tel que $n_1 > n_2$ où n_1 est l'indice du cœur et n_2 celui de la gaine.

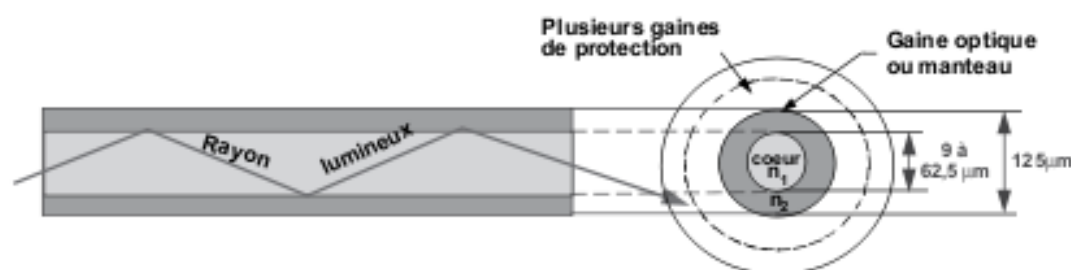


Figure 4.16 La fibre optique : guide de lumière.

Un système de transmission par fibre optique met en œuvre (figure 4.17) :

- ▶ un émetteur de lumière (transmetteur), constitué d'une diode électroluminescente **LED** (*Light Emitting Diode*) ou d'une diode **LASER** (*Light Amplification by Stimulated Emission of Radiation*), qui transforme les impulsions électriques en impulsions lumineuses ;
- ▶ un récepteur de lumière, constitué d'une photodiode de type **PIN** (*Positive Intrinsic Négative*) ou de type PDA (à effet d'avalanche) qui traduit les impulsions lumineuses en signaux électriques ;
- ▶ une fibre optique.

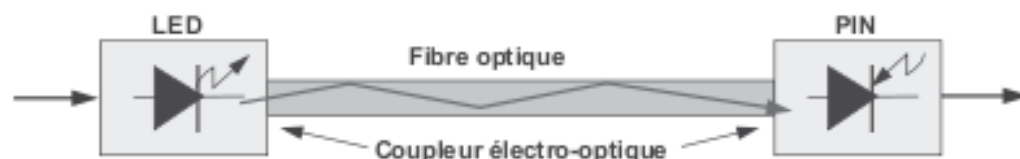


Figure 4.17 Principe d'une liaison optique.

■ Les différents types de fibres

Les fibres optiques sont distinguées en fonction du mode de propagation, de leur ouverture numérique et du mode de variation des indices entre le cœur et le manteau.

□ Les fibres à saut d'indice

Dans les fibres à saut d'indice (figure 4.18), le cœur d'indice n_1 est entouré d'une gaine d'indice n_2 . La variation d'indice entre le cœur et la gaine est brutale (saut d'indice). La propagation s'y fait par réflexion totale à l'interface cœur/gaine. Ce type de fibre admet plusieurs rayons qui se propagent sur des chemins différents ou modes de propagation. Ces différents trajets provoquent un étalement du signal (dispersion modale ou **DMD**, *Differential Mode Delay*), la fibre est alors dite multimode (**MMF**, *MultiMode optical Fiber*). La dispersion modale provoque un étalement du signal, ce qui limite la bande passante de la fibre et la distance franchissable.

En réduisant le diamètre du cœur, on réduit l'ouverture numérique (figure 4.19). Cette réduction, peut être telle que, pour une longueur d'onde donnée, la fibre n'admette plus qu'un seul rayon. La fibre est alors dite monomode (**SMF**, *Single Mode optical Fiber*). D'un coût plus élevé mais d'une bande passante beaucoup plus importante, elle est utilisée par les opérateurs de télécommunication pour réaliser des liaisons dites longues distances (MAN, WAN).

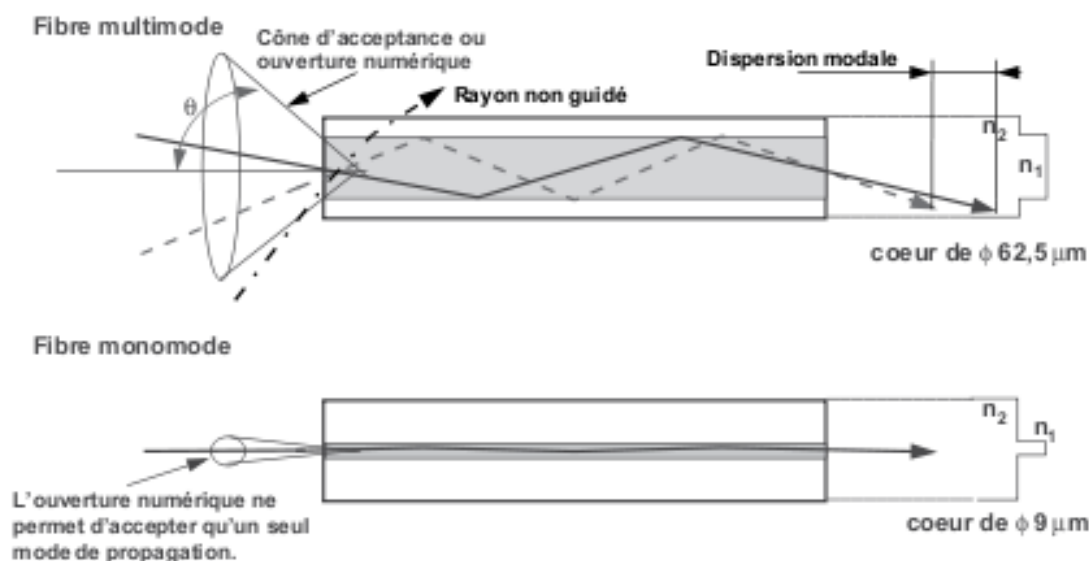


Figure 4.18 Les fibres à saut d'indice.

□ Les fibres à gradient d'indice

Un compromis a été trouvé avec les fibres à gradient d'indice (figure 4.19), l'indice du cœur décroît de façon continue, depuis le centre du cœur jusqu'à l'interface cœur/gaine suivant une loi parabolique. Dans ce type de fibre, tous les rayons sont focalisés au centre de la fibre, ils ont une trajectoire proche de la sinusoïde.

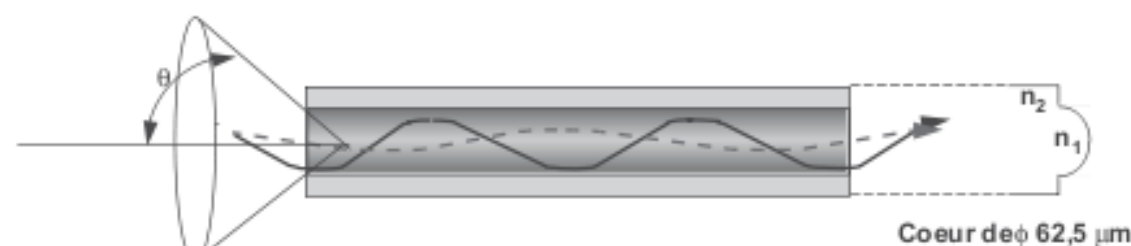


Figure 4.19 Les fibres à gradient d'indice.

■ La performance des fibres optiques

Dans une fibre optique, on montre que le produit bande passante par la distance est une constante. De ce fait, on exprime la bande passante par km. Compte tenu de la réponse en fréquence des fibres (figure 4.20) et des coupleurs optoélectroniques, on a défini trois plages d'utilisation appelées fenêtres optiques proches de l'infrarouge.

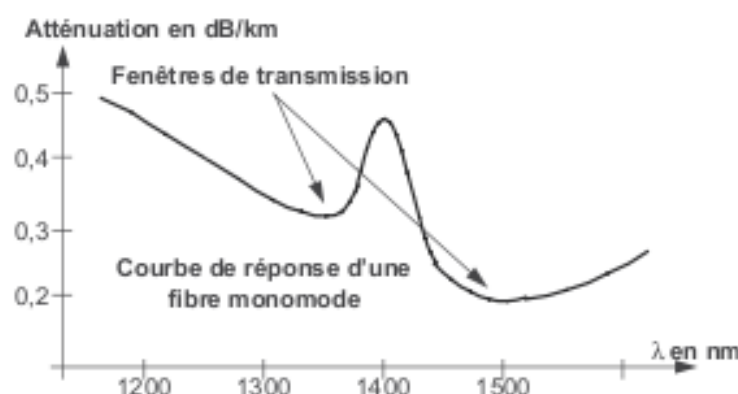


Figure 4.20 La notion de fenêtre optique.

La première fenêtre à 850 nm ($3,53 \cdot 10^5$ GHz) correspond à l'utilisation de coupleurs à coût minimal. Ce n'est pas l'optimum d'utilisation des fibres, mais dans des liaisons à faible distance, comme dans les réseaux locaux,

cette fenêtre est parfaitement adaptée. Généralement, on lui préfère la fenêtre de 1 300 nm ($2,3 \cdot 10^5$ GHz), l'atténuation n'est alors que d'environ 0,5 dB/km. La fenêtre située à 1 550 nm ($1,93 \cdot 10^5$ GHz) a l'avantage de ne présenter qu'une atténuation d'environ 0,2 dB/km, mais les coupleurs sont plus coûteux. Les performances des fibres optiques sont :

- ▶ bande passante importante,
- ▶ immunité électromagnétique,
- ▶ faible taux d'erreur 10^{-12} ,
- ▶ faible affaiblissement (0,2 à 0,5 dB/km),
- ▶ faible encombrement et poids,
- ▶ vitesse de propagation élevée (monomode),
- ▶ sécurité (absence de rayonnement à l'extérieur et difficulté de se mettre à l'écoute),
- ▶ légèreté.

Ces caractéristiques font des fibres optiques le support privilégié dans le domaine des télécommunications à haut débit et grande distance, dans les applications aéronautiques et navales (sous-marin) et dans les transmissions de données en milieu perturbé.

Si la pose de la fibre optique est aisée (pas de contraintes particulières), la connectique est assez délicate. Elle nécessite un outillage particulier et un savoir-faire certain. La figure 4.21 illustre les principaux connecteurs FO.



Figure 4.21 Les connecteurs optiques.

Les trois principaux types de connecteurs utilisés dans le domaine des entreprises sont les connecteurs ST, SC et LC :

- ▶ **ST**, d'origine AT&T (1985) est un connecteur carré en plastique à baïonnette de type « tourner/pousser ». Le couplage optique est réalisé par alignement de deux ferrures céramique ou métal. Le ST, standard *de facto*, est majoritairement utilisé pour les applications LAN multi-modes.
- ▶ **SC** ou *Subscriber Connector* est un connecteur rond à baïonnette de type encliquetable « push-pull » développé par NTT (1986). L'interface SC est présente dans un grand nombre d'équipements actifs (Ethernet, ATM, Sonet/SDH, Fiber Channel...);
- ▶ **LC** ou *Lucent Connector* est un connecteur à verrouillage à languette, surtout utilisé dans les équipements de la marque (Lucent et AVAYA).

■ La fibre optique dans les réseaux longue distance

L'UIT a spécifié (recommandation G.8723) une architecture de transport optique (OTN, *Optical Transport Networks*) comportant 3 niveaux (figure 4.22) :

- ▶ **OCH** (*Optical Channel layer*) mono longueur d'onde pour l'acheminement des données de bout en bout *via* un canal optique en assurant une certaine QoS ;
- ▶ **OMS** (*Optical Multiplex Section layer*) assure le multiplexage de longueur d'onde (supervision et gestion) ;
- ▶ **OTS** (*Optical Transmission Section layer*) assure la transmission d'un signal optique sur différent type de fibre (multimodes et nonomodes), il supervise le signal et assure l'amplification optique du signal (OA).

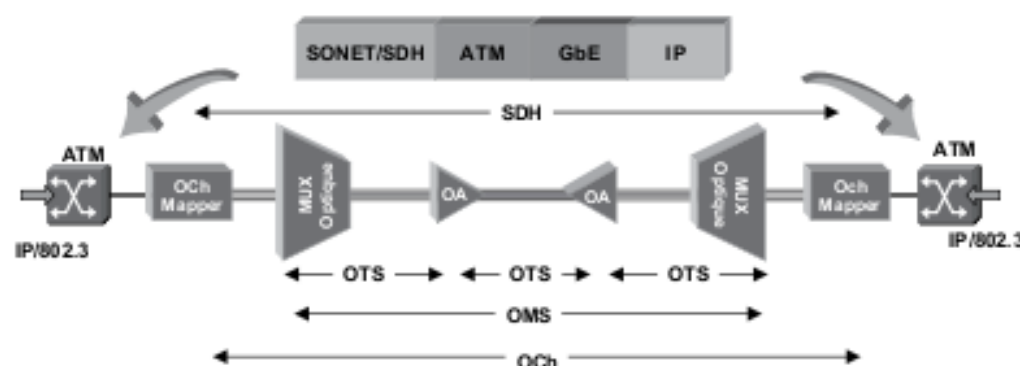


Figure 4.22 La transmission optique.

4.3 Les supports non guidés

4.3.1 Principe des liaisons hertziennes

Un conducteur rectiligne alimenté en courant haute fréquence ou radio-fréquence peut être assimilé à un circuit oscillant ouvert. Un tel circuit ou antenne d'émission rayonne une énergie, ou onde électromagnétique, qui résulte de la combinaison d'un champ magnétique et électrostatique. Cette énergie électromagnétique se propage sans support matériel. Lorsqu'elle est recueillie par un autre conducteur distant (antenne de réception), elle est transformée en un courant électrique similaire à celui d'excitation de l'antenne d'émission. La figure 4.23 illustre le principe d'une liaison radio-électrique. Les ondes électromagnétiques (OEM) ou ondes hertziennes¹ se propagent dans le vide à la vitesse de la lumière (c).

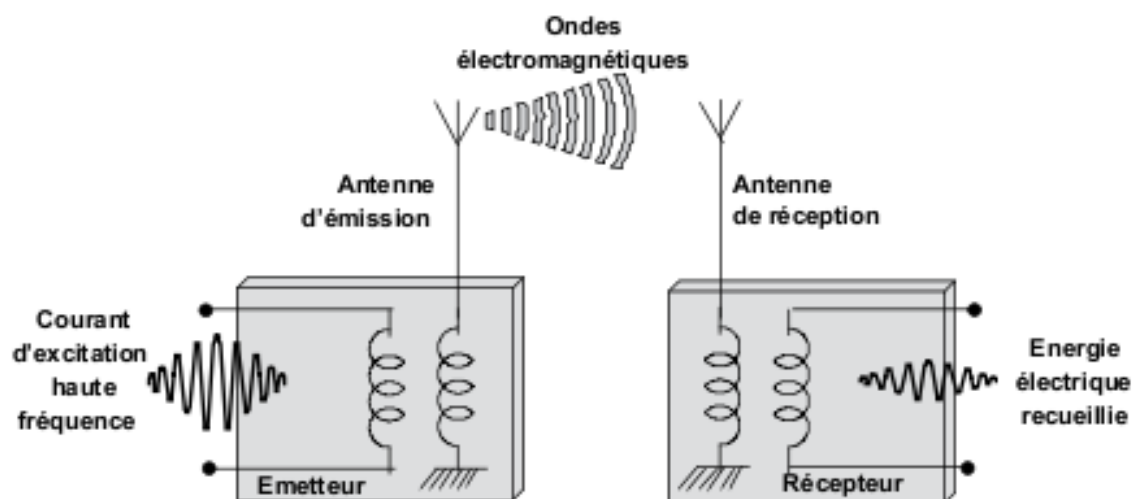


Figure 4.23 Principe d'une liaison radioélectrique.

Les ondes électromagnétiques ont une mise en œuvre aisée et sont d'un coût d'infrastructure généralement très faible devant les coûts de génie civil engendrés par le passage de câbles physiques. Les transmissions par ondes électromagnétiques sont utilisées chaque fois qu'il est nécessaire :

¹ Les ondes hertziennes ont été découvertes par le physicien allemand Heinrich Hertz en 1888.

- ▶ de diffuser une même information vers plusieurs utilisateurs (réseaux de diffusion),
- ▶ de mettre en relation des stations mobiles (réseaux de messagerie),
- ▶ de relier, à haut débit, deux entités éloignées (faisceaux hertziens) ou très éloignées (satellites de communication),
- ▶ de permettre la mobilité comme pour les réseaux locaux sans fil (WiFi) ou la téléphonie mobile.

Les ondes électromagnétiques subissent peu d'affaiblissement dans l'atmosphère, sauf par temps de brouillard ou de pluie où les particules d'eau absorbent l'énergie des ondes. Selon la longueur d'onde, certaines matières absorbent toute l'énergie de l'onde, créant ainsi de véritables zones d'ombre. D'autres matériaux réfléchissent les ondes, ainsi un récepteur peut recevoir plusieurs fois l'information, une fois par le trajet direct et une ou plusieurs fois par des trajets réfléchis (figure 4.24).

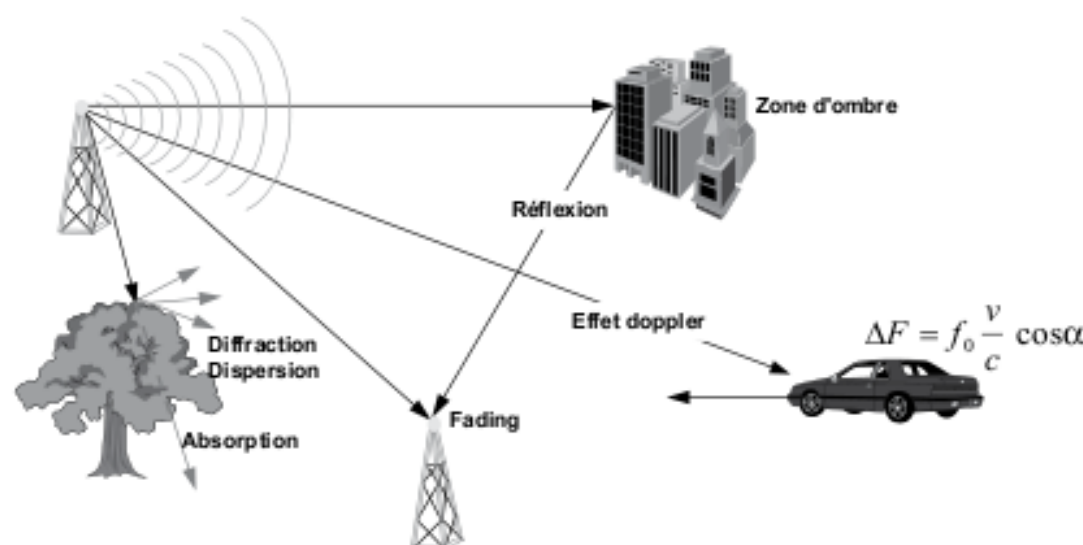


Figure 4.24 Synthèse des phénomènes agissant sur la transmission des OEM.

Chaque type de liaison ou d'application utilise des bandes de fréquences différentes. L'espace de fréquences utilisables est limité. On appelle **canal radio** ou **canal de transmission** la bande de fréquences réservée à une communication. La figure 4.25 décrit le spectre de fréquences. Les ondes radioélectriques s'étendent de quelques dizaines de kilohertz (ondes longues ou grandes ondes) à plus du térahertz (ondes quasi optiques).

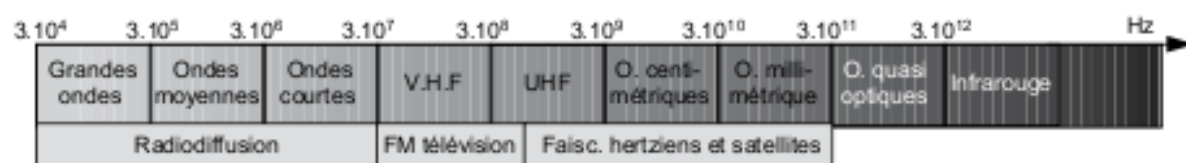


Figure 4.25 Le spectre des fréquences.

L'usage des fréquences n'est pas libre, celui-ci est réglementé. Au niveau international, les fréquences sont gérées par l'UIT-TS (Union internationale des télécommunications – *Telecommunication Standardization*, ex-CCITT et CCIR). Les différents domaines d'utilisation se voient attribuer une bande de fréquences, elle-même divisée en canaux. L'attribution locale des fréquences est généralement le fait d'organismes nationaux : l'ANF (Agence nationale des fréquences) et ARCEP (Autorité de régulation des communications électroniques et des postes) en France.

4.3.2 Les faisceaux hertziens

Les ondes radioélectriques peuvent, dans certains cas, remplacer avantageusement les liaisons filaires (cuivre ou optique). Pour diminuer les puissances d'émission, la technique des faisceaux hertziens utilise des antennes très directives. L'antenne réelle est placée au foyer optique d'une parabole qui réfléchit les ondes en un faisceau d'ondes parallèles très concentré, limitant ainsi la dispersion de l'énergie radioélectrique. En réception, l'antenne est aussi placée au foyer optique de la parabole. Tous les rayons reçus parallèlement à l'axe optique de la parabole sont réfléchis vers le foyer optique, on recueille ainsi un maximum d'énergie (figure 4.26).

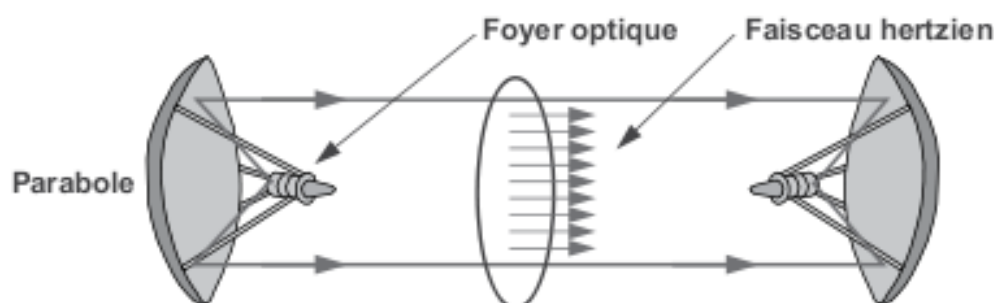


Figure 4.26 Principe des faisceaux hertziens.

Les liaisons infrarouges et lasers constituent un cas particulier des liaisons hertziennes. Elles sont généralement utilisées pour réaliser des liaisons privées de l'ordre de quelques centaines de mètres afin d'interconnecter des réseaux privés. Pratiquement abandonnés dans les liaisons longues distances, les faisceaux hertziens ont retrouvé un regain d'intérêt dans les liaisons entre relais de téléphonie mobile.

4.3.3 Les liaisons satellitaires

La nécessité de disposer de stations relais rend difficile la réalisation de liaisons hertziennes à très grande distance, notamment pour les liaisons transocéaniques. C'est pourquoi, dès les années 1960, on s'est orienté vers l'utilisation de satellites relais. Mais, ce n'est qu'avec l'apparition de porteurs capables de satelliser sur des orbites d'environ 36 000 km qu'il a été possible de réaliser des liaisons permanentes. À cette altitude, ces satellites ont une vitesse angulaire de rotation identique à celle de la Terre, ce qui correspond à une période de révolution de 23 h 56'. Ces satellites paraissent fixes à un observateur terrestre (satellite géostationnaire ou **géosynchrone**).

■ Principe

Une station terrestre émet vers le satellite un flux d'information (voie montante). Le satellite n'est qu'un simple répéteur, il régénère les signaux reçus, réalise une transposition de fréquence et les réémet en direction de la Terre (voie descendante). La figure 4.27 illustre le principe d'une liaison satellitaire.

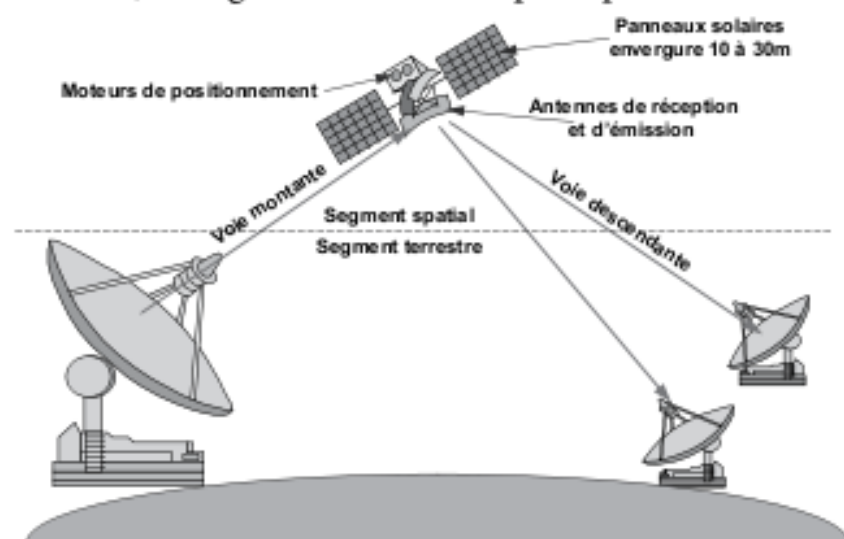


Figure 4.27 Principe d'une liaison satellitaire.

■ Les différents types de systèmes satellitaires

Compte tenu des temps de propagation des satellites géostationnaires, des systèmes à orbites plus basses ont été définis. Selon leur orbite, les systèmes satellitaires sont regroupés en trois familles. On distingue : les orbites stationnaires (GEO), moyennes (MEO) et basses (LEO). Le tableau 4.4 résume les caractéristiques de ces satellites.

Tableau 4.4 Synthèse des caractéristiques des différents systèmes de satellites.

	GEO Geostationary Earth Orbit	MEO Medium Earth Orbit	LEO Low Earth Orbit
Altitude	36 000 km	2 000 à 12 000 km	800 à 2 000 km
Type d'orbite	Circulaire	Elliptique ou circulaire	Elliptique ou circulaire
Plan de rotation	Équatorial	Quelconque	Quelconque
Temps de transmission	240 ms	110 à 150 ms	Environ 50 ms
Terre-satellite			
Permanence spatiale et temporelle (spatiale : communiquer en tout point ; temporelle : en un point à tout moment)	OUI Trois satellites couvrent la Terre (sauf les pôles)	NON (orbite défilante) Constellation de satellites	NON (orbite défilante) Constellation de satellites
Applications	Téléphonie fixe, télévision, transmission de données	Téléphonie mobile, transmission de données	Téléphonie mobile, transmission de données
Débit	Jusqu'à 155 Mbit/s	De 9,6 à 38 kbit/s	De 2,4 kbit/s à 155 Mbit/s

5

Les modes de transmission

5.1 L'organisation des échanges

5.1.1 Liaison simplex, half-duplex, full duplex

Les transmissions entre deux systèmes peuvent être unidirectionnelles (l'échange n'a lieu que dans une seule direction), on parle alors de **liaison simplex**, (figure 5.1). Si les correspondants peuvent, alternativement, remplir les fonctions d'émetteur et de récepteur, la liaison est dite : **liaison à l'alternat** ou **half duplex**. Lorsque l'échange peut s'effectuer simultanément dans les deux sens, la liaison est appelée **bidirectionnelle intégrale** ou **full duplex**.

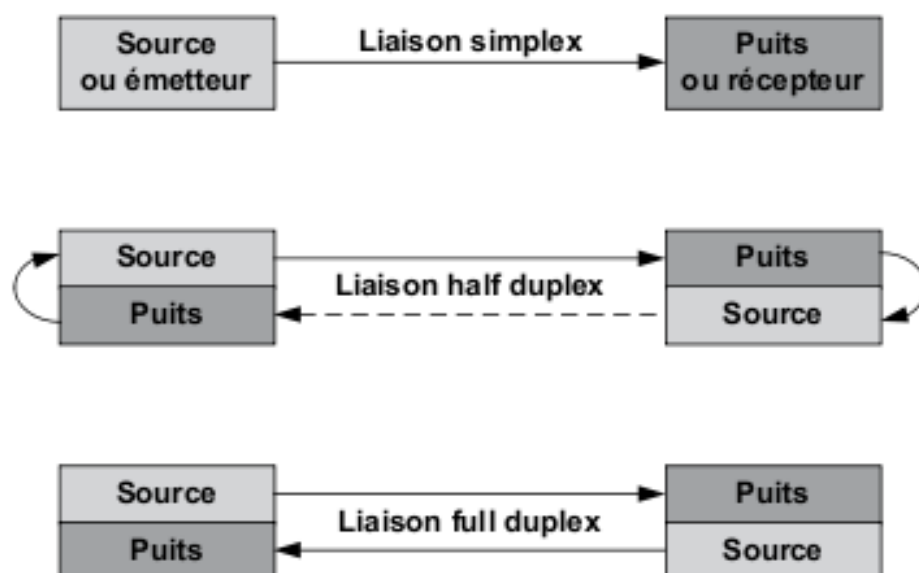


Figure 5.1 L'organisation fonctionnelle des échanges.

5.1.2 Transmission série et parallèle

L'information élémentaire à transmettre est le mot (8, 16, 32 ... n bits). En interne, les calculateurs transfèrent les données *via* un bus : un fil par bit (retour commun). Un bus transmet simultanément tous les bits d'un même mot machine, la transmission est dite transmission parallèle. La communication entre machines peut se réaliser de même. Cependant, la transmission parallèle soulève de nombreux problèmes techniques. Pour des distances importantes, on lui préfère la transmission série : les bits sont transmis successivement sur un support unique.

Si on désigne par **temps bit** le temps d'émission d'un bit sur le support, et en considérant que ce temps est identique pour la transmission parallèle et série de la figure 5.2, on constate qu'il faut seulement trois temps bit pour transmettre le mot « ISO » en transmission parallèle, alors que la transmission série nécessite huit temps bit pour transmettre la seule lettre « O ».

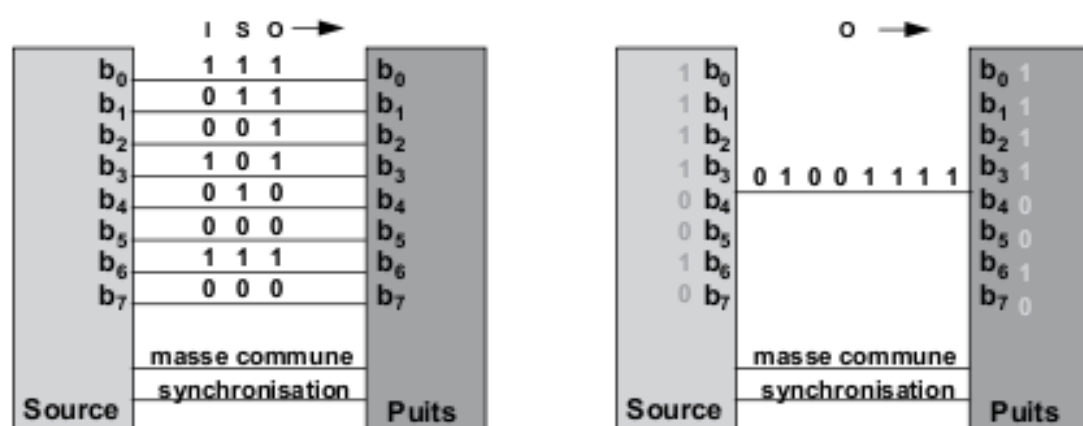


Figure 5.2 Transmission parallèle, transmission série.

5.1.3 Transmission asynchrone, transmission synchrone

■ Notion d'horloge

Les bits sont émis sur la ligne à une certaine cadence. Cette cadence est définie par une horloge dite **horloge émission**. Pour décoder correcte-

ment la suite de bits reçue, le récepteur doit examiner ce qui lui arrive à la même cadence que celle de l'émission des bits sur le support. L'horloge du récepteur et celle de l'émetteur doivent « battre » en harmonie (fréquence et phase). À cette fin, l'horloge du récepteur est asservie sur celle de l'émetteur, cette fonction se nomme synchronisation des horloges (figure 5.3).

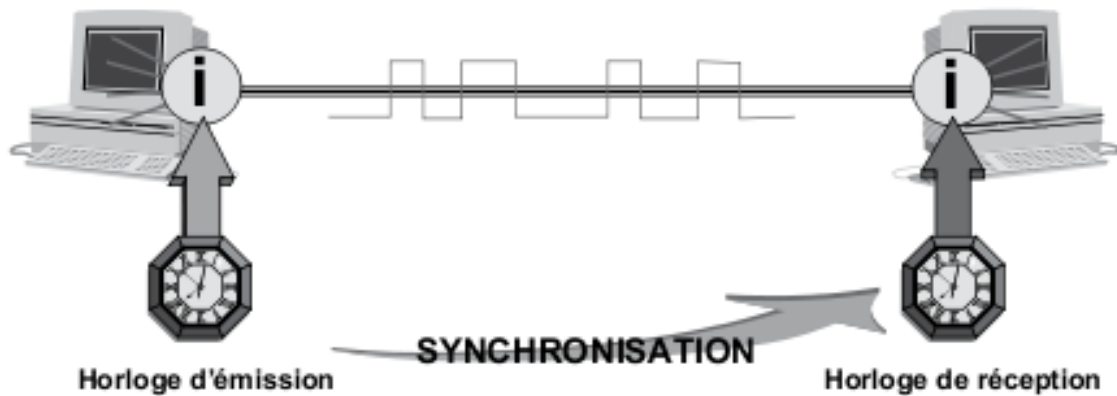


Figure 5.3 Principe de la synchronisation.

Selon le mode de synchronisation de l'horloge du récepteur sur celle de l'émetteur, on distingue deux types de transmission : les transmissions asynchrones et les transmissions synchrones.

Dans les transmissions asynchrones on synchronise l'horloge en début d'émission, ensuite, l'horloge du récepteur bat librement, la transmission ne peut excéder le caractère (temps de stabilité acceptable de l'horloge récepteur), les horloges sont indépendantes. Dans les transmissions synchrones, on maintient en permanence une relation de phase stricte entre les horloges émission et réception. Le signal de synchronisation est déterminé à partir des données transmises. Cette « extraction » de l'horloge est généralement remplie par le DCE (ETCD) plus connu sous le terme de modem, même si ce terme est parfois improprement utilisé (figure 5.4).



Figure 5.4 Principe de la synchronisation d'une liaison informatique.

■ Transmission asynchrone

En transmission asynchrone, les caractères émis sont précédés d'un signal de synchronisation : le bit de start. Entre chaque caractère, pour garantir la détection du bit de start suivant, la ligne est remise à l'état zéro. Ce temps de repos minimal varie d'un à deux temps bit, il constitue le **bit de stop** (figure 5.5).

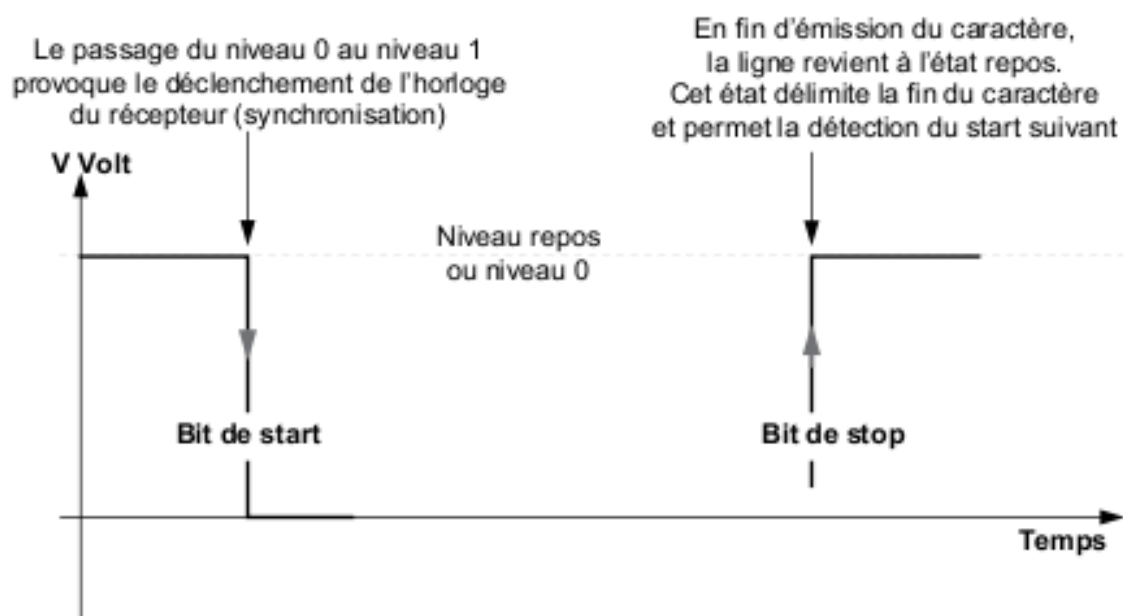


Figure 5.5 Principe de la synchronisation en transmission asynchrone.

Les transmissions asynchrones s'effectuent selon un ensemble de règles régissant les échanges (protocole). On distingue deux types de protocoles asynchrones (figure 5.6)

- ▶ Les **protocoles en mode caractère** : la transmission a lieu caractère par caractère. L'intervalle de temps qui sépare chaque caractère peut être quelconque (multiple de la fréquence d'horloge).
- ▶ Les **protocoles en mode bloc** : les caractères sont rassemblés en blocs. L'intervalle de temps entre l'émission de deux blocs successifs peut être quelconque (multiple de la fréquence d'horloge).

Asynchrone en mode caractère



Asynchrone en mode bloc



Figure 5.6 Mode caractère et mode bloc.

■ Transmission synchrone

En transmission synchrone, la synchronisation des horloges émission et réception est maintenue en permanence entre l'émetteur et le récepteur. Il est alors possible de transmettre des blocs de taille importante. Pour maintenir la synchronisation durant les silences de l'émetteur, des signaux spécifiques sont émis (bourrage de ligne).

Si la synchronisation de lecture des bits est maintenue par ce procédé (synchronisation bit), il n'en est pas de même de la frontière entre chaque caractère. Celle-ci est délimitée en transmission asynchrone par les caractères de start et de stop, ces caractères sont absents de la transmission synchrone. Alors, afin de cadrer la lecture des bits sur une frontière d'octet, chaque bloc transmis est précédé d'une séquence de synchronisation spécifique dite synchronisation caractère. La synchronisation caractère est réalisée par reconnaissance dans le train binaire d'une séquence binaire prédéfinie. Lorsque cette séquence est reconnue par le récepteur, celui-ci se cale alors sur une lecture de 8 bits en 8 bits du flot de bits pour lire un flux d'octets (figure 5.7). Cette séquence binaire ou **fanion** sert aussi de délimiteur de début et de fin de bloc. En l'absence d'émission, les fanions servent aussi au bourrage de ligne (maintien de la synchronisation bit).

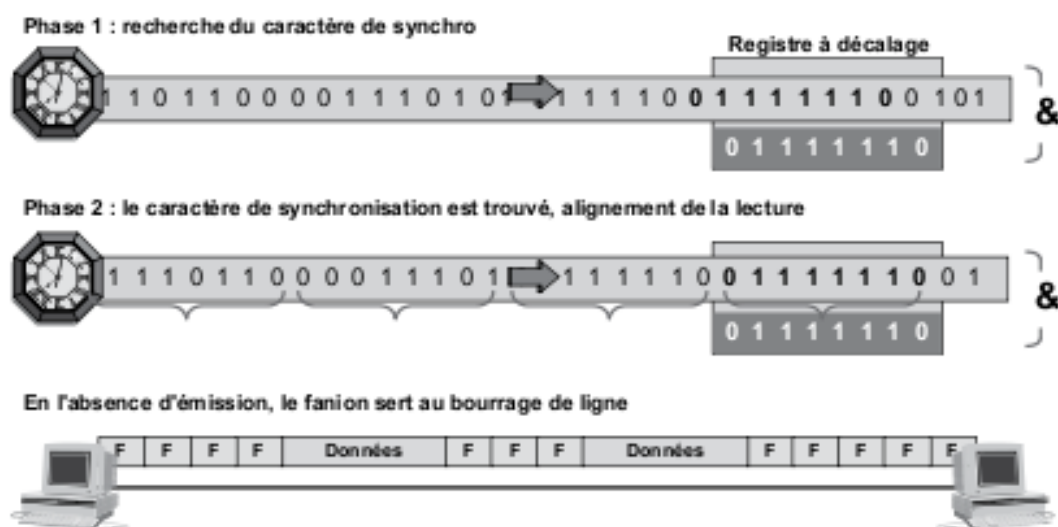


Figure 5.7 Principe de la synchronisation caractère.

5.2 Transmissions bande de base et large bande

La bande passante limitée du système de transmission conduit à modifier le signal électrique pour l'adapter aux contraintes physiques du système de transmission. Deux types d'adaptation ou techniques de transmission sont envisageables :

- ▶ La première consiste à modifier légèrement le signal, elle est essentiellement destinée à réduire la composante continue. Cependant, les composantes hautes fréquences étant fortement atténuées, la transmission sera limitée en distance : c'est la transmission en bande de base.
- ▶ La seconde translate le spectre du signal à émettre dans une bande de fréquences mieux admise par le système, c'est la transmission large bande.

5.2.1 La transmission en bande de base

■ Définitions

On qualifie de systèmes de transmission en bande de base les systèmes qui n'introduisent pas d'écart de fréquence entre les signaux émis et ceux

reçus. Cette définition n'exclut nullement des modifications du signal pour mieux l'adapter aux caractéristiques du support de transmission.

On appelle **codage**, l'opération qui fait correspondre à un symbole appartenant à un alphabet, une représentation binaire (codage à la source : ASCII...). On désigne par **transcodage**, ou **codage en ligne**, l'opération qui consiste à substituer au signal numérique (représentation binaire) un signal électrique mieux adapté à la transmission. Cette transformation est réalisée par un codeur/décodeur appelé Émetteur/récepteur en bande de base (**ERBdB**) souvent improprement appelé modem bande de base.

■ Les fonctions d'un codeur/décodeur en bande de base

Le signal numérique, issu du calculateur, présente une composante continue¹ non nulle. Cette composante continue est inutile, elle ne transporte aucune information et provoque un échauffement (effet Joule) des organes d'extrémité (transformateurs d'isolement). Le comportement de filtre passe-bas du système introduit une distorsion de phase qui provoque l'étalement du signal. L'absence de transition, lors de la transmission d'une longue suite de 0 ou de 1, introduit un risque de perte de synchronisation des horloges. Ces différentes considérations conduisent à :

- ▶ transformer le signal numérique en un autre, tel que la composante continue soit réduite à son minimum voire supprimée ;
- ▶ choisir une méthode de codage pour que le spectre du nouveau signal soit mieux adapté aux caractéristiques du support de transmission ;
- ▶ et enfin, pour maintenir la synchronisation, assurer un minimum de transitions, même lors de la transmission de longues séquences de 1 ou de 0.

En résumé, le transcodage, ou codage en ligne, a essentiellement pour objet de supprimer la composante continue, d'adapter le spectre au canal de transmission et de maintenir la synchronisation de l'horloge de réception. On distingue trois types de code :

¹ La composante continue représente la « valeur moyenne » du signal pour un intervalle de temps donné.

- ▶ ceux qui effectuent un codage des 1 et 0 (Manchester...);
- ▶ ceux qui ne codent que les 1 ou les 0 (bipolaire...);
- ▶ ceux qui substituent à un ensemble de n bits un autre ensemble de m bits (nBmB).

Exemple de codage : le code Manchester

En symétrisant le signal par rapport au potentiel de référence (0 volt), on diminue la composante continue. Pour cela, on représente les 1 (ou les 0) par une valeur $+V$ et les 0 (ou les 1) par $-V$. Ce codage élémentaire connu sous le nom de code **NRZ** (*No Return to Zero*, non-retour à zéro) est à la base de tous les codes (figure 5.8). Cependant, le spectre de ce signal est relativement large ; il présente un maximum de puissance à la fréquence zéro, ce qui correspond à une composante continue encore importante.

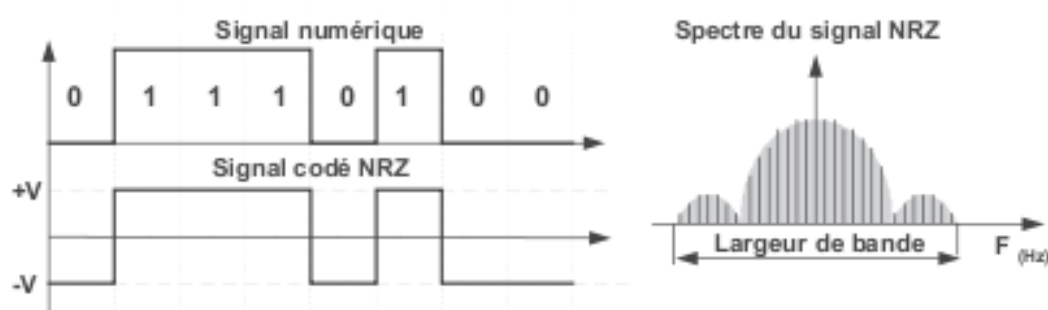


Figure 5.8 Le codage NRZ.

Le codage NRZ symétrise la valeur 1 et la valeur 0 par rapport à un niveau potentiel nul. Cependant, ce codage a encore une composante continue non nulle et ne présente aucune transition lors de longues séquences de 0 ou de 1.

Avec une transition au milieu de chaque temps bit, le codage Manchester (figure 5.9) remédie à l'absence d'information de synchronisation. La transition est croissante pour les 0, décroissante pour les 1. Le sens des transitions est significatif, ce qui pose des problèmes en cas d'inversion des fils de liaison. Multipliant les transitions, le codage Manchester a un spectre très large, il est utilisé dans les réseaux locaux de type Ethernet à 10 Mbit/s.

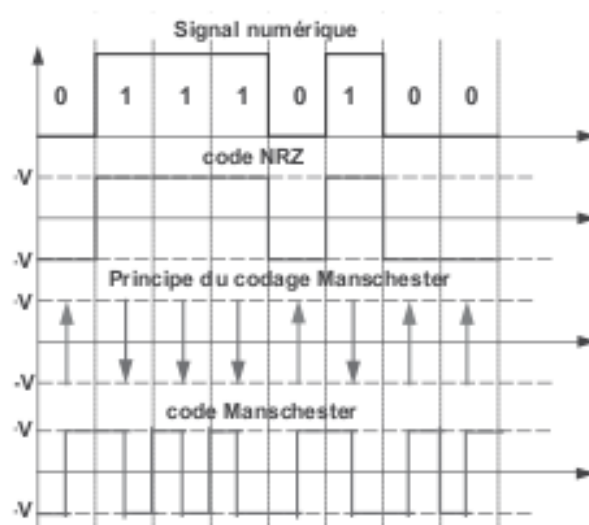


Figure 5.9 La construction du code Manchester.

5

■ Limitation de la transmission en bande de base

□ Critère de Nyquist

La transmission en bande de base est une technique simple à mettre en œuvre, mais elle est limitée par la bande passante du canal de communication et par le rapport signal sur bruit de celui-ci.

Il existe une relation étroite entre le nombre maximal de symboles (impulsions électriques) que le système peut admettre et la bande passante en fréquence de celui-ci. Aussi, le nombre maximal d'impulsions électriques et la bande passante sont liés par la relation appelée **critère de Nyquist** :

$$R_{\max} \leq 2 \cdot BP$$

où $R_{\max}$ désigne le nombre maximal de transitions qu'un système peut supporter, et est appelé **rapidité de modulation**. La rapidité de modulation, grandeur analogue à une fréquence, s'exprime en baud et représente le nombre d'instant élémentaires du signal par unité de temps, on l'appelle aussi vitesse de signalisation.

Cependant ce n'est qu'une limite « électronique ». Imaginons que, durant un temps élémentaire, le symbole prenne plusieurs états (figure 5.10), la quan-

tité d'informations transportée alors par un symbole est supérieure à 1 bit. En fait, débit binaire et rapidité de modulation sont liés par la relation :

$$D = R \cdot Q = R \cdot \log_2 (1/p)$$

avec D , débit binaire exprimé en bit/s
 R , rapidité de modulation en bauds
 Q , quantité d'informations en bits
 p , probabilité d'apparition d'un état.

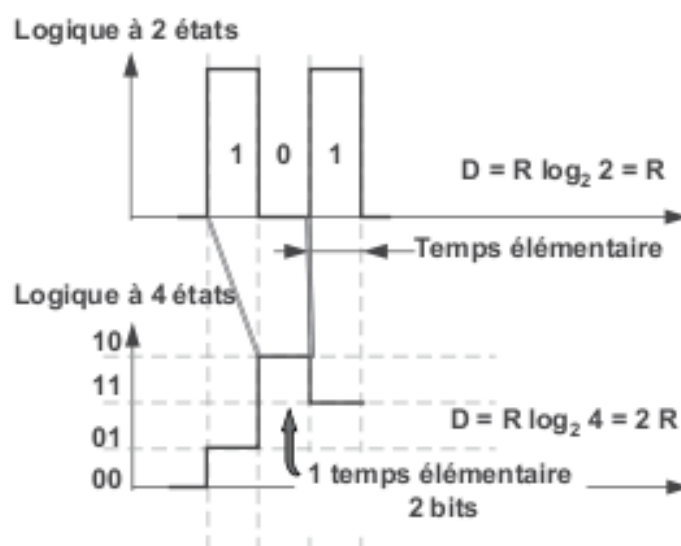


Figure 5.10 Notion de valence du signal.

Si on appelle valence du signal (v) le nombre d'états que peut prendre le signal durant un temps élémentaire ($v = 1/p$), le débit s'exprime, alors, par la relation :

$$D = R \cdot \log_2 v = 2 \cdot BP \cdot \log_2 v$$

avec D , débit binaire en bit/s
 v , valence du signal, valant $1/p$
 R , rapidité de modulation

Dans le cas de l'exemple précédent, la rapidité de modulation est égale au débit binaire. Si durant un temps élémentaire, le signal peut prendre plusieurs

valeurs, par exemple 4 (figure 5.10), la probabilité d'apparition d'une valeur est de 0,25. Dans ces conditions, le débit du canal est :

$$D = R \cdot \log_2 (1/0,25) = R \cdot \log_2 4 = 2 \cdot R \text{ bit/s}$$

Le débit binaire est le double de la rapidité de modulation. C'est ainsi qu'il est possible d'augmenter, sur un canal de transmission de bande passante limitée, le débit binaire. L'opération qui consiste à faire correspondre à un ensemble de symboles binaires (00, 01...000, 001...) un signal électrique spécifique représentatif de cette combinaison est réalisée par un codeur (codage en ligne).

On ne peut augmenter indéfiniment le nombre d'états du signal (valence), car les niveaux d'amplitude à discriminer deviennent si faibles qu'ils ne peuvent être distingués du bruit (figure 5.11).

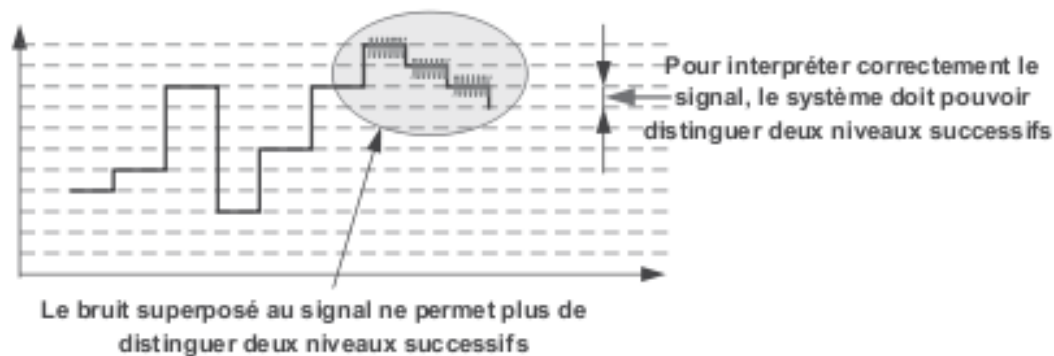


Figure 5.11 La limitation du nombre d'états par le bruit.

Reprenant les travaux de Nyquist, Claude Shannon a montré qu'en milieu perturbé, le nombre maximal d'états discernables, ou **valence**, est donné par la relation où S représente le niveau du signal et N celui du bruit (*Noise*) :

$$n = \sqrt{1 + S/N}$$

La capacité maximale de transmission d'un canal est alors de :

$$C = 2 \cdot BP \cdot \log_2 n = BP \cdot \log_2 [1 + (S/N)]$$

Application au réseau téléphonique commuté (RTC)

Quelle est la capacité maximale de transmission sur une voie RTC (Réseau téléphonique commuté) caractérisée par une bande passante de 300 / 3 400 Hz et un rapport signal sur bruit de 1 000 ?

La rapidité de modulation maximale de ce canal est :

$$R_{\max} = 2 \cdot BP = 2 (3\,400 - 300) = 6\,200 \text{ bauds}$$

La capacité de transmission est donnée par la relation de Shannon :

$$C = BP \cdot \log_2 [1 + (S/N)]$$

$$C = (3\,400 - 300) \log_2 (1 + 1\,000) \approx 3\,100 \cdot 3,32 \log_{10}(1\,000)$$

$$C = 3\,100 \cdot 3,32 \cdot 3$$

$$C = 30\,876 \text{ bits/s}$$

Ce débit maximal théorique correspond aux performances maximales que l'on peut obtenir sur une ligne téléphonique¹.

5.3 La transmission large bande

En transmission large bande, le spectre du signal numérique est translaté autour d'une fréquence centrale appelée **porteuse**. La translation de spectre résout les deux problèmes posés par la transmission en bande de base : dispersion du spectre (étalement du signal) et la monopolisation du support qui interdit le multiplexage. Elle est réalisée par un organe appelé **modulateur**. En réception le signal doit subir une transformation inverse, il est démodulé. Le modem, contraction de **modulation/démodulation**, est un équipement qui réalise la modulation des signaux en émission et leur démodulation en réception (figure 5.12).

¹ Le débit maximal sur ligne téléphonique ordinaire (BP = 300 – 3 400 Hz) est aujourd'hui atteint par les modems V.34 bis (33 600 bit/s).

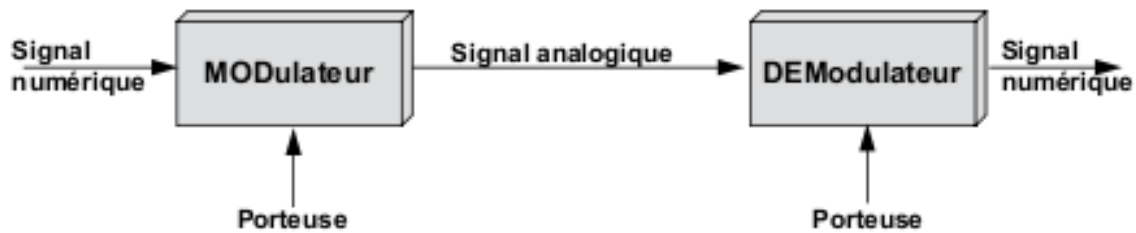


Figure 5.12 La transmission analogique.

La dégradation du signal impulsionnel de la bande de base est rapide, la distance franchissable est limitée à quelques kilomètres. Le signal sinusoïdal est plus résistant, d'où l'idée de substituer au signal impulsionnel, un signal sinusoïdal et de modifier l'un de ses paramètres en fonction du signal numérique d'origine : c'est la **modulation**. Un signal sinusoïdal est de la forme :

$$U_{(t)} = A_0 \sin(\omega_0 t + \varphi_0) \text{ avec } \omega_0 = 2\pi f_0$$

Sur ce signal, on peut faire varier (figure 5.13) :

- ▶ l'amplitude A_0 , c'est la modulation d'amplitude (**ASK**, *Amplitude Shift Keying*) ;
- ▶ la fréquence f_0 , c'est la modulation de fréquence (**FSK**, *Frequency Shift Keying*) ;
- ▶ la phase φ_0 , c'est la modulation de phase (**PSK**, *Phase Shift Keying*).

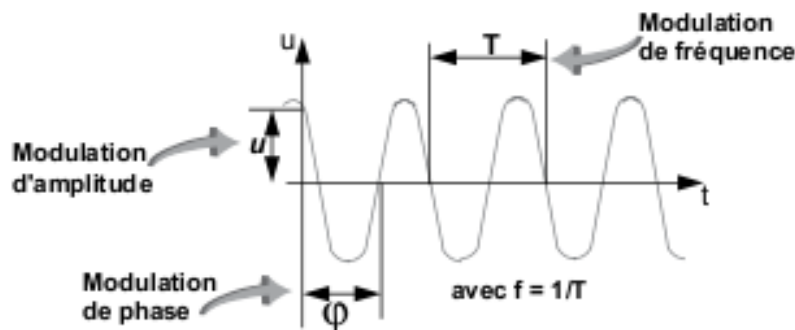


Figure 5.13 Le signal sinusoïdal et la modulation.

Les modems utilisent aujourd'hui une combinaison de la modulation d'amplitude et de phase. On obtient des schémas de modulation complexes mais très

efficaces. Ce type de modulation appelé modulation en amplitude à porteuse en quadrature (**MAQ**, ou **QAM Quadrature Amplitude Modulation**) résiste bien au bruit et autorise des débits élevés avec une rapidité de modulation relativement faible. La figure 5.14 gauche représente le diagramme spatial d'un schéma de modulation à 16 états (MAQ16). Cette technique est limitée par l'erreur de phase introduite par le bruit (figure 5.14 droite).

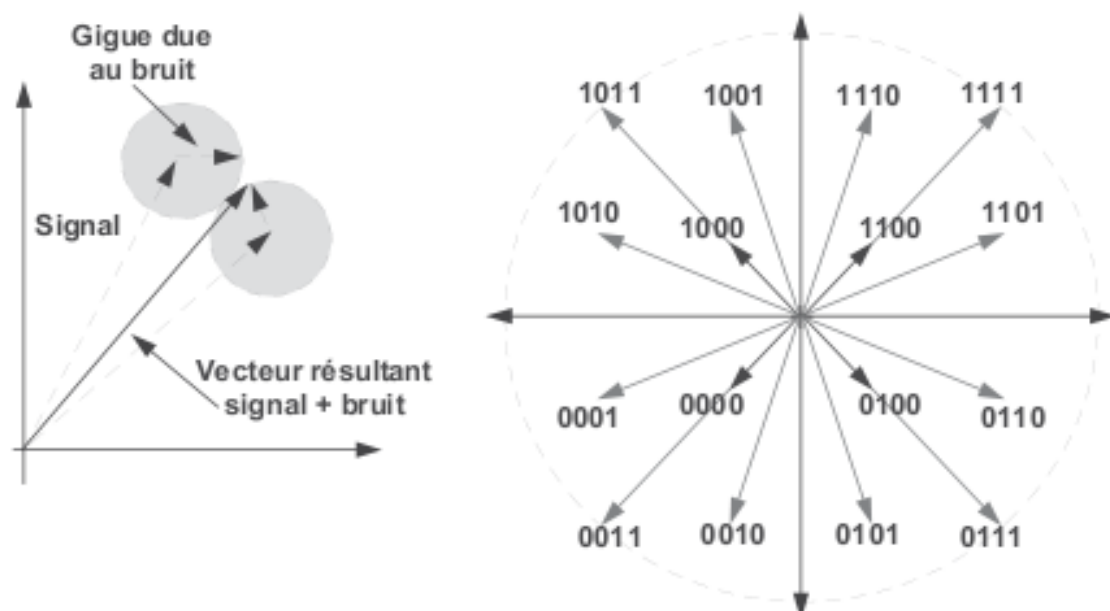


Figure 5.14 Principe de la modulation MAQ (MAQ16) et ses limitations.

Les modems de dernière génération peuvent mettre en œuvre des codages jusqu'à 64 états, autorisant ainsi des débits élevés avec une rapidité de modulation faible.

6

Le multiplexage

Lors de la réalisation d'une liaison de transmission de données, le responsable réseau et télécoms d'une entreprise doit rechercher la meilleure solution en termes d'efficacité et de coût. Cet objectif d'optimisation des moyens de transmission conduit, lorsque les utilisateurs ne monopolisent pas la ressource de manière permanente, au partage de celle-ci, c'est le **multiplexage**.

Le **multiplexeur** est un équipement qui met en relation un utilisateur avec un autre par l'intermédiaire d'un support partagé par plusieurs utilisateurs. Un multiplexeur de n voies simule, sur une seule ligne, n liaisons point à point. Chaque voie d'entrée est dénommée voie incidente, le support partagé voie composite (figure 6.1).

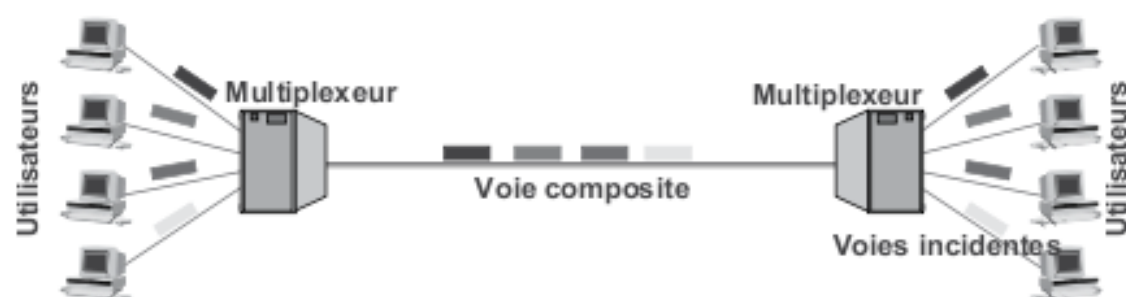


Figure 6.1 Principe du multiplexage.

L'opération de regroupement des voies incidentes sur un même support s'appelle le multiplexage. Le démultiplexage consiste à restituer à chaque destinataire les données qui lui sont destinées. Un MUX (abréviation utilisée pour désigner un multiplexeur) est un système symétrique, il comporte à la fois un organe de multiplexage et un organe de démultiplexage (liaison *full duplex*).

Le partage de la voie composite peut être celui :

- ▶ de la bande disponible, chaque voie dispose en permanence d'une fraction de la bande passante, c'est le **multiplexage fréquentiel** ou **spatial** ;
- ▶ du temps d'utilisation de la voie, chaque voie utilise durant un temps pré-déterminé toute la bande disponible, c'est le **multiplexage temporel**.

6.1 Le multiplexage spatial

Le multiplexage fréquentiel (**FDM**, *Frequency Division Multiplexing*) correspond à une juxtaposition fréquentielle de voies et à une superposition des signaux dans le temps. La bande passante du support est divisée en canaux (voies). Chaque voie est modulée (transposition de fréquence) par une porteuse différente. Le démultiplexage correspond, à l'aide de filtres, à la séparation de chacune des voies (voies spatiales) puis à la démodulation de chacun des signaux. La figure 6.2 illustre le principe d'un multiplexeur fréquentiel.

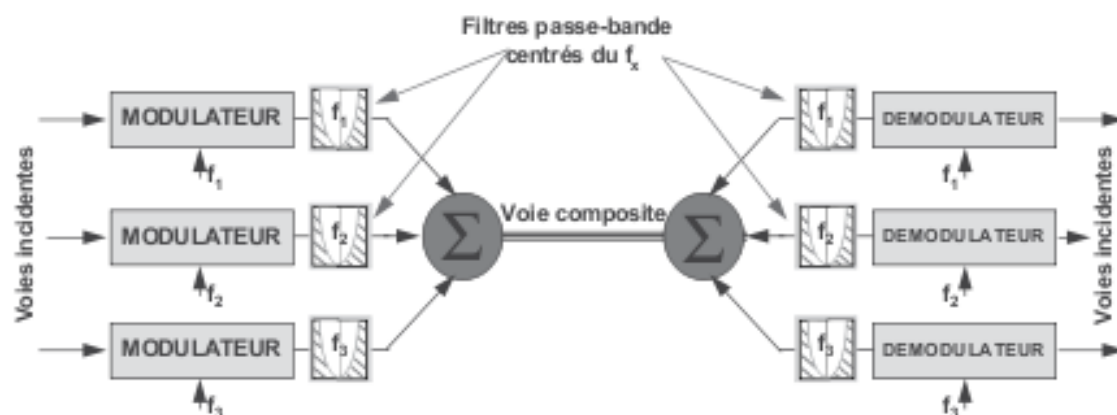


Figure 6.2 Principe du multiplexage fréquentiel.

Avant l'apparition des techniques de numérisation, le multiplexage fréquentiel a été utilisé pour constituer les premiers réseaux de téléphonie.

6.1.1 Le multiplexage de longueur d'onde

Les liaisons optiques mettent en œuvre un cas particulier du multiplexage fréquentiel (figure 6.3) : le multiplexage de longueur d'onde (**WDM**, *Wavelength Division Multiplexing*). Le WDM consiste à injecter dans une même fibre optique plusieurs trains de signaux numériques, chacun sur une longueur d'onde distincte.

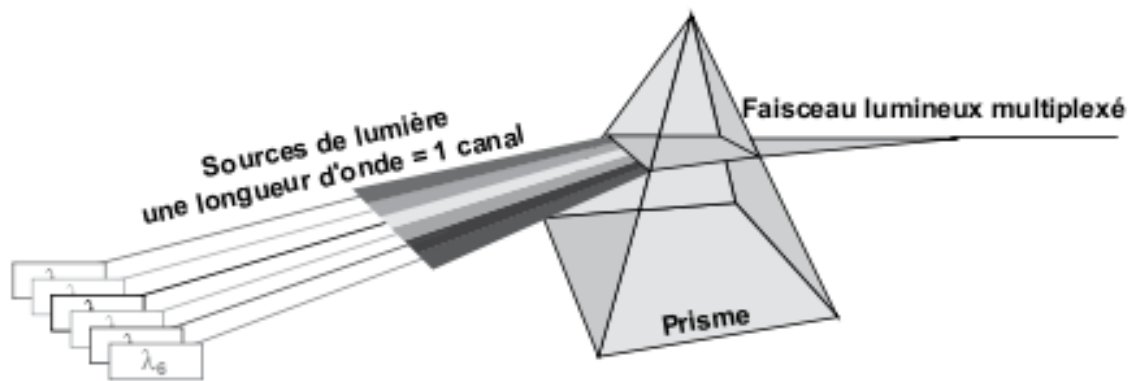


Figure 6.3 Principe du multiplexage de longueur d'onde.

Les technologies sont dites *Dense WDM* (**D-WDM**) lorsque l'espacement entre canaux est égal ou inférieur à 100 GHz. Les réalisations actuelles autorisent le multiplexage de 4 (1997), 16 et 32 (1998), 64 (1999) et aujourd'hui 105 longueurs d'onde avec des trains numériques incidents de 2,5, 10 voire 40 Gbit/s.



Figure 6.4 Principe d'une liaison WDM.

La figure 6.4 illustre un système de transmission en multiplexage de longueurs d'onde. Le faisceau lumineux est régénéré par un amplificateur optique à fibre dopée à l'erbium (**EDFA**, *Erbium Doped Fibre Amplifier*). Les pas de régénération sont d'environ 100 km pour les liaisons terrestres

et 50 à 80 km pour les liaisons océaniques. Le tableau 6.1 décrit les différents modes de multiplexage de longueur d'onde.

Tableau 6.1 Les différents modes de WDM.

Mode	Appellation	Espacement	Débit par λ	Nb. de λ	Débit potentiel
C-WDM	Coarse WDM	1,6 – 0,8 nm	1.25 & 2,5	8 à 16	10G à 25G
WDM	WDM	0,6 nm	10 & 40	32	320G à 1,28T
D-WDM	Dense WDM	0,4 - 0,2 nm	10 & 40	300	3T à 12T
U-DWDM	Ultra Dense WDM	0,08 nm	10 & 40	1022	10T à 40T

6.2 Le multiplexage temporel

Quand les utilisateurs n'utilisent pas en permanence le support, il existe des espaces de temps (silences) qui peuvent être utilisés par d'autres utilisateurs. Les multiplexeurs temporels relient une voie incidente d'entrée à une voie incidente de sortie durant un intervalle de temps prédéterminé. Cet intervalle de temps ou IT, réservé à un couple émetteur/récepteur, constitue une voie temporelle (figure 6.5).

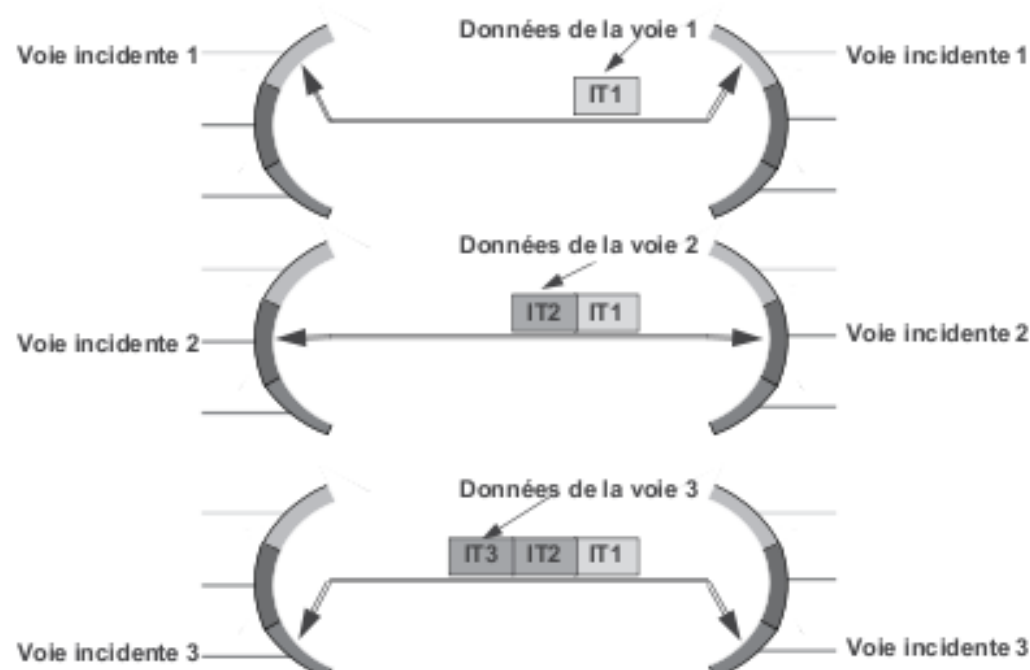


Figure 6.5 Principe du multiplexage temporel.

Dans le premier schéma de la figure 6.5, le couple de multiplexeurs met en relation les utilisateurs raccordés aux voies identifiées « voies incidentes 1 » ; l'intervalle de temps suivant, les utilisateurs raccordés aux voies 2, puis ceux raccordés aux voies 3. Un tel système transporte des bits, le multiplexeur n'interprète pas les données qu'il transporte, il est dit transparent au protocole. L'arrivée des données est indépendante du fonctionnement du multiplexeur. Les informations qui arrivent, sur une voie, pendant la période de scrutation des autres voies incidentes sont mémorisées (figure 6.6). Les multiplexeurs nécessitent de la mémoire et introduisent un retard de transmission qui peut être important vis-à-vis du temps de transfert sur le support.

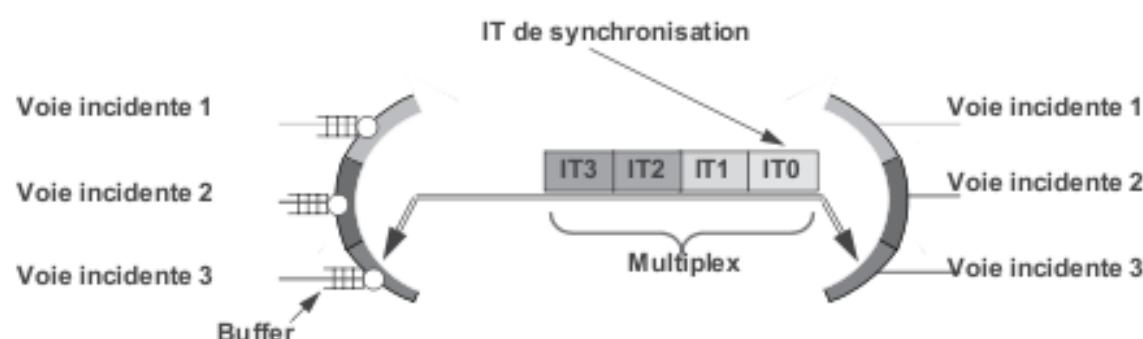


Figure 6.6 Structure élémentaire d'un multiplexeur.

6.3 Le multiplexage inverse

Le multiplexage inverse (**IM**, *Inverse Multiplexing*) consiste en l'agrégation de plusieurs liens bas débit pour obtenir un débit, vu de l'utilisateur, égal à la somme des débits des liens agrégés (figure 6.7). Cette technique est mise en œuvre par les opérateurs pour accroître le débit de leurs liens optiques (WDM).

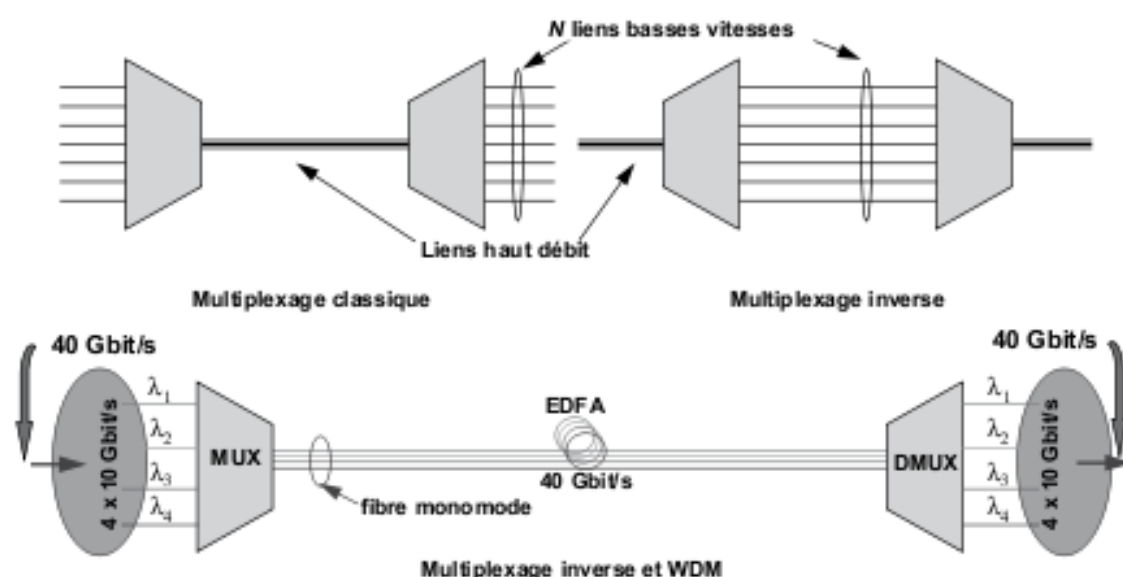


Figure 6.7 Principe du multiplexage inverse.

6.4 Conclusion

La limitation des débits est essentiellement due aux caractéristiques physiques des supports, mais les limites sont loin d'être atteintes. Les progrès des techniques de traitement du signal et de celles du codage des informations permettent d'améliorer la résistance au bruit et laissent espérer des débits se rapprochant de plus en plus des limites physiques. Dans l'attente de l'évolution de ces techniques, le multiplexage de longueur d'onde est la solution adoptée par tous les opérateurs pour accroître, à moindre coût, le débit de leurs liens.

3

Les protocoles de liaison

7

Les fonctions élémentaires

7.1 Notion de protocole

Dans les chapitres précédents, nous avons étudié les mécanismes à mettre en œuvre pour transmettre un flot de bits entre deux systèmes distants. Cependant, il ne suffit pas de lire correctement les bits reçus, encore faut-il les présenter pour qu'ils puissent être traduits en données utilisables par les applications, c'est le rôle des **protocoles**. On appelle protocole un ensemble de conventions préétablies pour réaliser un échange de données entre deux entités.

Les protocoles se déclinent selon deux familles. Ceux qui « garantissent » un transfert fiable de données en garantissant la délivrance de celles-ci, ce sont les protocoles dits **orientés connexion**, les blocs de données se nomment alors « trames », et ceux, beaucoup plus simples, qui se contentent d'envoyer des données sans en vérifier la bonne délivrance : ils sont dits en **mode sans connexion** ou mode datagramme ; dans ce dernier cas, l'unité de données est nommée « **datagramme** ». Cependant certaines fonctions sont communes aux deux types, c'est le cas de la délimitation des données et éventuellement du contrôle d'erreur...

7.2 La délimitation des données

7.2.1 Notion de fanion

À l'instar des transmissions asynchrones où les bits de start et de stop encadrent les bits d'information, en transmission synchrone un caractère spécial ou une combinaison particulière de bits, le **fanion**, permet de repérer le début et la fin du bloc de données transmis (figure 7.1).



Figure 7.1 La délimitation des données par fanions.

Le fanion assure trois fonctions essentielles :

- ▶ il délimite les données ;
- ▶ en l'absence de données à émettre, il permet de maintenir la synchronisation de l'horloge réception (fréquence et phase) ;
- ▶ en identifiant le fanion, le récepteur peut se caler correctement sur une frontière d'octets (**synchronisation caractère**) et, par conséquent, traduire le flux de bits reçus en un flux d'octets.

7.2.2 Notion de transparence

L'utilisation d'un caractère spécifique pour indiquer le début ou la fin d'un bloc de données interdit l'usage de ce caractère dans le champ des données. En conséquence, si on veut transmettre, en tant que données, le caractère ou une combinaison binaire similaire à celle du fanion, il faut prévoir un mécanisme particulier dit **mécanisme de transparence au caractère**, si le fanion est un caractère, ou **mécanisme de transparence binaire**, si le fanion est une combinaison de bits.

Le mécanisme de transparence consiste à « baliser » le caractère à protéger par un autre caractère dit **caractère d'échappement**. Ce caractère inséré à l'émission devant le caractère à protéger (le faux fanion) doit lui-même être protégé s'il apparaît dans le champ de données (figure 7.2).

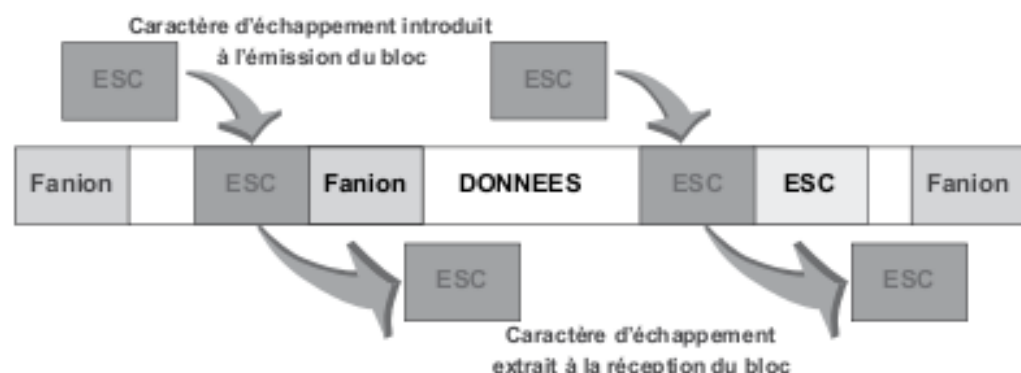


Figure 7.2 Principe de la transparence au caractère.

L'émetteur insère le caractère d'échappement devant le caractère à protéger. En réception, l'automate examine chaque caractère pour découvrir le fanion de fin. S'il rencontre le caractère d'échappement, il l'élimine et n'interprète pas le caractère qui le suit, il le délivre au système comme un caractère ordinaire.

Dans d'autres protocoles, un champ particulier est réservé aux informations de contrôle. Ce champ peut contenir une combinaison binaire quelconque. Ces protocoles sont dits orientés bits. Dans ces protocoles, le fanion est constitué de la combinaison binaire « 01111110 » soit 0x7E. La transparence binaire est assurée par l'insertion d'un « 0 » tous les 5 bits à « 1 » consécutifs. Ainsi, seul, le fanion contient une combinaison binaire de plus de 5 bits à 1 consécutifs (01111110). Cette technique de transparence dite du **bit de bourrage** (*Bit stuffing*) est illustrée figure 7.3.

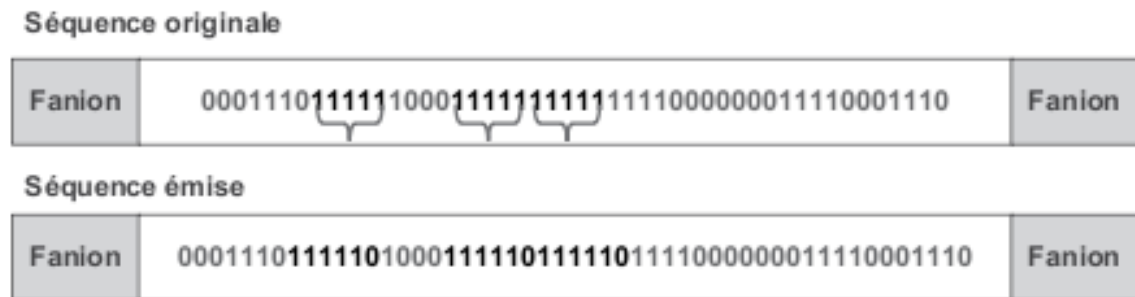


Figure 7.3 La technique du bit de bourrage.

7.3 Le contrôle d'erreur

Le contrôle d'erreur consiste à s'assurer que les données reçues n'ont pas été altérées durant la transmission. On appelle **détection d'erreur** le mécanisme mis en œuvre par le système destinataire pour vérifier la validité des données reçues. La détection d'erreur repose sur l'ajout au bloc de données à transmettre une information supplémentaire (clé) déduite des informations transmises (figure 7.4).



Figure 7.4 Principe de la correction d'erreur par redondance d'information.

7.3.1 Technique dite du bit de parité

La technique du **bit de parité** consiste à ajouter à la séquence binaire à protéger un bit, telle que la somme des bits à 1 transmis soit paire (bit de parité) ou impaire (bit d'imparité). Cette arithmétique modulo 2 est simple, mais elle n'introduit qu'une faible redondance. La protection apportée est limitée à quelques bits : le caractère. La figure 7.5 illustre le mécanisme de calcul du bit de parité.

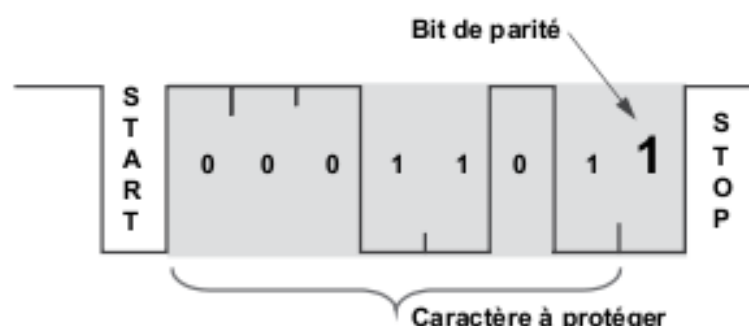


Figure 7.5 Le contrôle de parité dans les transmissions asynchrones.

Dans les transmissions synchrones, les caractères sont envoyés en blocs (figure 7.6). La technique du bit de parité est insuffisante, elle est complétée d'une autre information : le **LRC** (*Longitudinal Redundancy Check*).

Caractère à transmettre	Bit de parité	Caractère à transmettre	Bit de parité	...	Caractère à transmettre	Bit de parité
-------------------------	---------------	-------------------------	---------------	-----	-------------------------	---------------

Figure 7.6 La structure d'un bloc de caractères protégé par LRC.

Dans ce mode de contrôle, dit de parité à deux dimensions, un caractère : le LRC est ajouté au bloc transmis (figure 7.6). Chaque bit du caractère LRC correspond à la « parité » des bits de même rang que lui

de chacun des caractères composant le message : le premier bit du LRC est la parité de tous les premiers bits de chaque caractère, le second de tous les deuxièmes bits... Le caractère ainsi constitué est ajouté au message.

■ Détection d'erreur par clé calculée

Dans les systèmes à clé calculée, une séquence de contrôle (CTL1) déduite d'une opération mathématique appliquée au message à émettre est envoyée avec le message. Le récepteur effectue la même opération. Si le résultat trouvé (CTL2) est identique à la clé calculée par la source (CTL1), le bloc est réputé exact, dans le cas contraire le bloc est rejeté (figure 7.7).

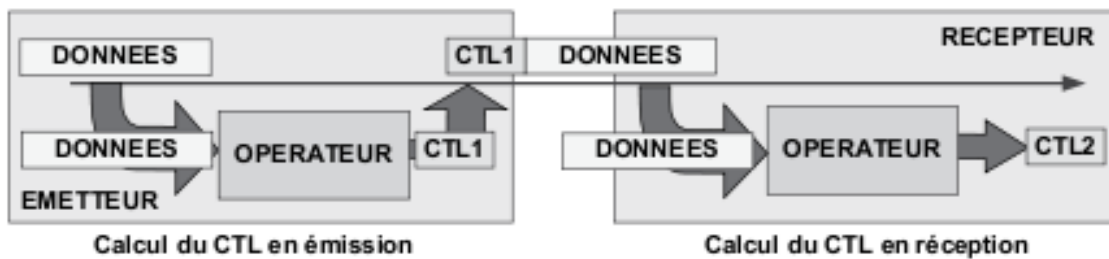


Figure 7.7 Principe de la détection d'erreur par clé calculée.

Dans la pile de protocoles TCP/IP utilisé par Internet, le mode de calcul du mot de contrôle se rapproche des techniques de parité. Le mot de contrôle sur 16 bits ou total de contrôle est le complément à 1 de la somme en complément à 1 des mots de 16 bits composant le message.

■ Les codes cycliques ou détection par clé calculée

Dans la détection par clé calculée, l'information redondante, la clé (CRC, *Cyclic Redundancy Check*), est déterminée par une opération mathématique complexe appliquée au bloc de données à transmettre et est transmise avec celui-ci (figure 7.8).

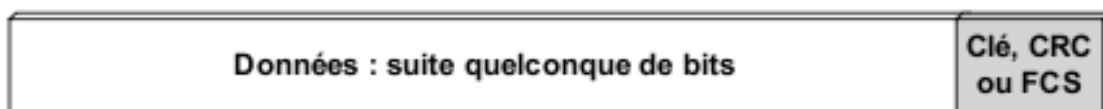


Figure 7.8 La structure d'un bloc de bits protégé par clé calculée.

La méthode de contrôle par clé calculée considère le bloc de N bits à transmettre comme un polynôme de degré $N - 1$: $P_{(x)}$. Ce polynôme est divisé par un autre, dit polynôme générateur $G_{(x)}$, selon les règles de l'arithmétique booléenne ou arithmétique modulo 2. Le reste de cette division $R_{(x)}$ constitue le CRC parfois appelé aussi **FCS** (*Frame Check Sequence*). Le CRC calculé est transmis à la suite du bloc de données. En réception, le destinataire effectue la même opération sur l'intégralité (données et CRC) du bloc reçu (figure 7.9), on montre, qu'en absence d'erreur le reste de la division est nulle.

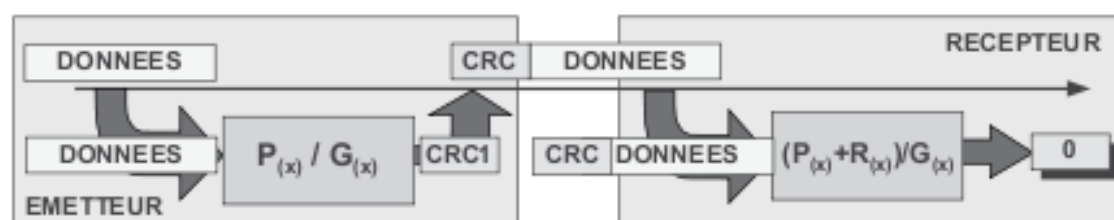


Figure 7.9 La détection d'erreur par CRC.

8

Exemples de protocole de liaison

8.1 SLIP (Serial Line Internet Protocol)

Historiquement premier protocole de liaison de la pile TCP/IP, SLIP est un protocole sans connexion. SLIP (RFC 1055) est un protocole asynchrone orienté bloc. Très simple, il n'effectue que la délimitation de trames. En fait, SLIP se contente d'émettre les données (datagramme) et signale au récepteur la fin du message par l'émission d'un caractère de délimitation (END, 0xC0). N'indiquant que la fin d'un datagramme, SLIP est susceptible d'introduire une confusion entre le bruit de ligne et les données effectives. Aussi, afin d'assurer une délimitation correcte du datagramme, la plupart des implémentations fait précéder l'émission de données du caractère « END ».



Figure 8.1 Le format de la trame SLIP.

Dans la figure 8.1, le bruit de ligne est assimilé à des données par le récepteur. L'émission d'un caractère END, avant l'émission du datagramme IP, marque la fin du « datagramme bruit » qui sera ultérieurement éliminé par les couches supérieures, car considéré comme un bloc erroné de données. Le datagramme réel est donc bien délimité.

Le caractère END (192D ou 0xC0) est utilisé comme délimiteur de début et de fin. SLIP assure la transparence en remplaçant le caractère END

par la séquence ESC_SLIP¹, ESC_END soit 0xDB, 0xDC. Le caractère d'échappement (ESC_SLIP) est lui-même protégé par l'insertion de la séquence ESC_SLIP, ESC_ESC soit les valeurs 0xDB, 0xDD. La figure 8.2 illustre ce mécanisme.

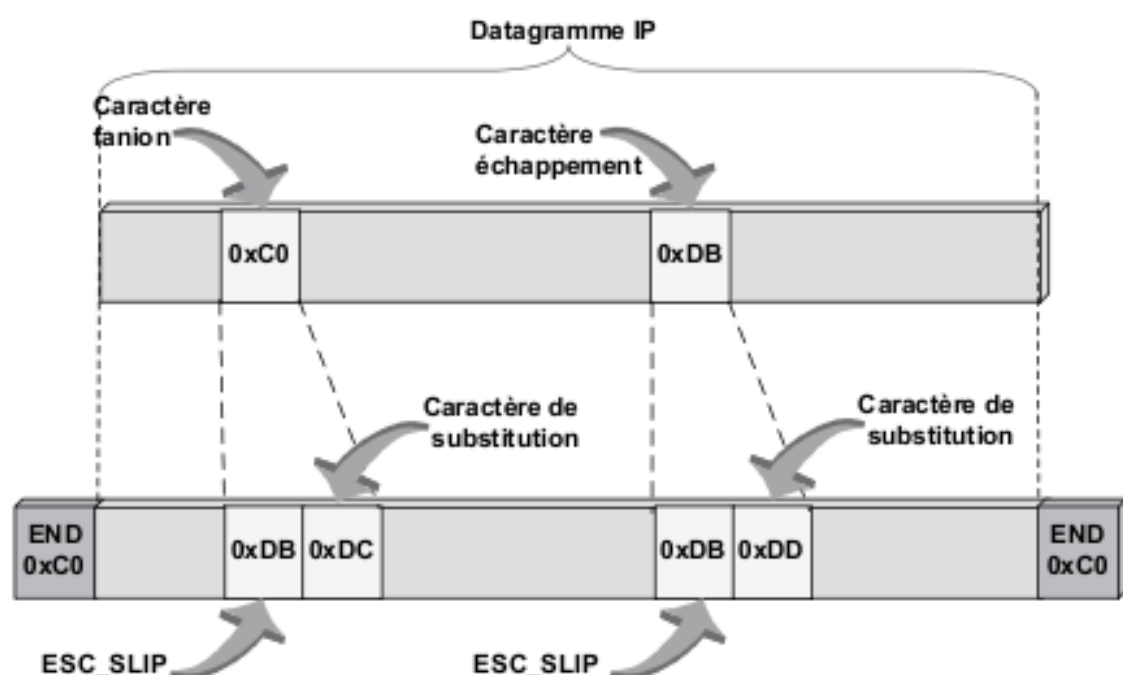


Figure 8.2 Le mécanisme de la transparence dans SLIP.

Peu fiable, SLIP n'est pratiquement plus que dans des liaisons locales entre un mini-ordinateur et ses terminaux ; on lui préfère les protocoles en mode connecté comme HDL ou PPP.

8.2 HDLC, High Level Data Link Control

HDLC est un protocole ligne dit de **point à point**. Dérivé de SDLC (*Synchronous Data Link Control*) d'IBM, normalisé par le CCITT

¹ Le caractère d'échappement utilisé par le protocole SLIP est différent du caractère d'échappement du code ASCII. Pour le distinguer, nous le noterons ESC_SLIP, de même celui utilisé par le protocole PPP sera noté ESC_PPP.

(UIT-T) en 1974 et l'ISO (1976), HDLC a été à l'origine de nombreux autres protocoles de ligne (figure 8.3).

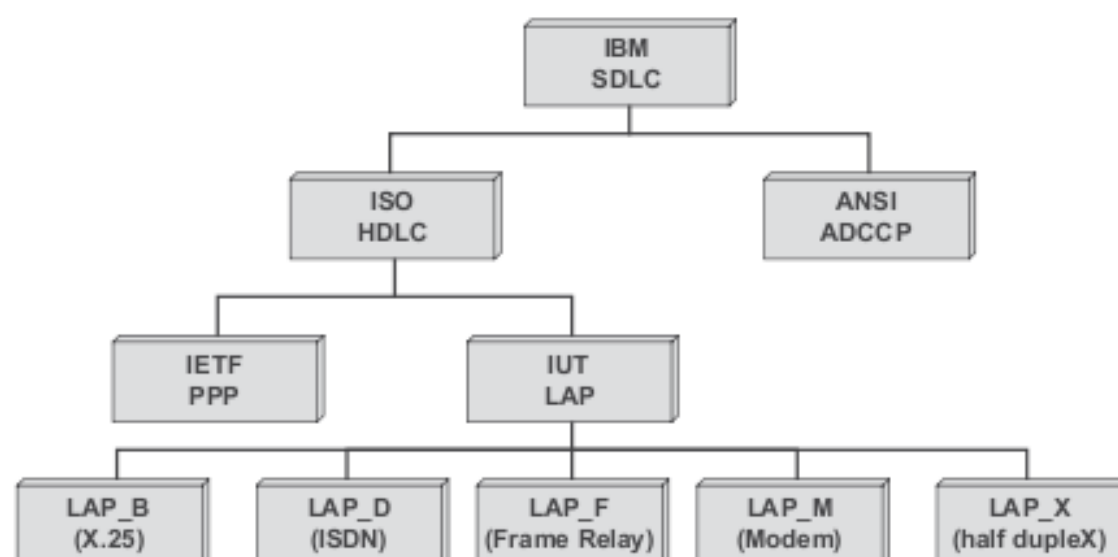


Figure 8.3 Généalogie d'HDLC.

L'unité de transfert d'HDLC est la trame (*frame*), chaque trame est délimitée par un caractère spécifique : le fanion ou *flag* aussi employé pour maintenir, en l'absence de données à transmettre, la synchronisation entre les trames. La figure 8.4 schématise une liaison HDLC, les symboles « F » représentent les fanions envoyés durant les silences pour maintenir la synchronisation.

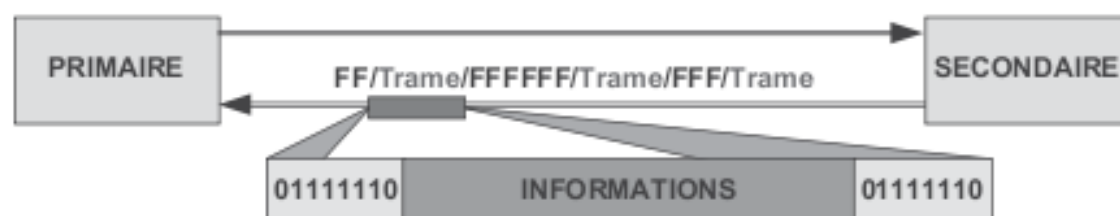


Figure 8.4 La liaison HDLC.

HDLC est un protocole qui utilise un mode de signalisation (ensemble des informations de gestion de l'échange) dans la bande. À cet effet, on distingue trois types de trames (figure 8.5).



Figure 8.5 Les fonctions et trames correspondantes d'HDLC.

Les trames d'information ou trames **I** assurent le transfert de données ; les trames de supervision ou trames **S** (*Supervisor*) en effectuent le contrôle (accusé de réception...), enfin les trames non numérotées ou trames **U** (*Unnumbered*) supervisent la liaison. Les trames **U** sont, en principe, des trames de signalisation.

8.2.1 Les mécanismes d'HDLC

■ La numérotation des trames

Le principe de base de toute transmission en mode orienté connexion repose sur l'envoi (*Send*) d'un bloc d'information. L'émetteur s'arrête alors (*Stop*) dans l'attente (*Wait*) d'un accusé de réception. À la réception de l'acquittement, noté **ACK** pour *Acknowledge*, l'émetteur envoie le bloc suivant (figure 8.6 gauche).

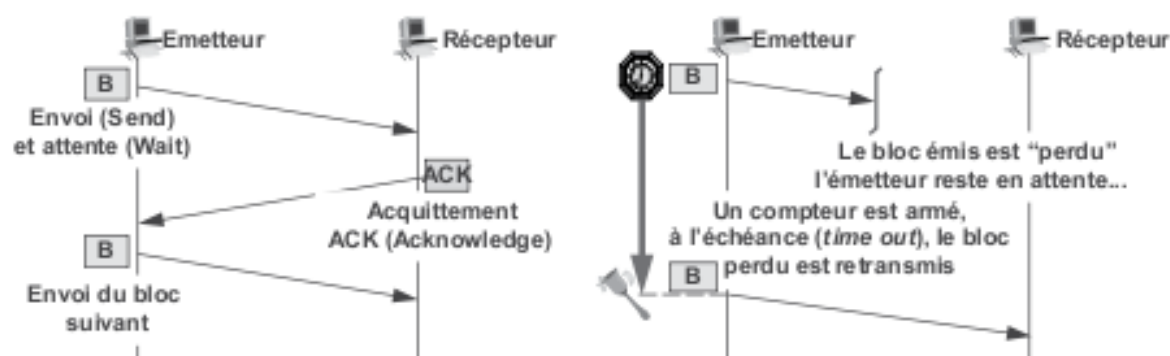


Figure 8.6 Le mode Send et Wait et la reprise sur temporisation.

En cas d'erreur de transmission, le bloc reçu est rejeté. Le bloc est dit perdu, il n'est pas acquitté. L'émetteur reste alors en attente. Pour évi-

ter un blocage de la transmission, à l'émission de chaque bloc de données, l'émetteur arme un temporisateur (*Timer*). À l'échéance du temps imparti (*Time Out*), si aucun accusé de réception (ACK) n'a été reçu, l'émetteur retransmet le bloc non acquitté, cette technique porte le nom de reprise sur temporisation (**RTO**, *Retransmission Time Out*) ou correction d'erreur sur temporisation (figure 8.6 droite).

Une difficulté survient si la perte concerne l'ACK. En effet, bien que les données aient été correctement reçues, l'émetteur les retransmet sur temporisation. Les informations sont ainsi reçues deux fois. Pour éviter la duplication des données et pour associer l'acquiescement à la bonne trame, il est nécessaire d'identifier les blocs. À cet effet, l'émetteur et le récepteur entretiennent des compteurs, compteurs qui sont transportés dans l'en-tête protocolaire (figure 8.7) :

- ▶ N_s , numéro du bloc envoyé ;
- ▶ N_r , numéro du bloc attendu, ce numéro acquitte toutes les trames reçues jusqu'à (N_r-1) .

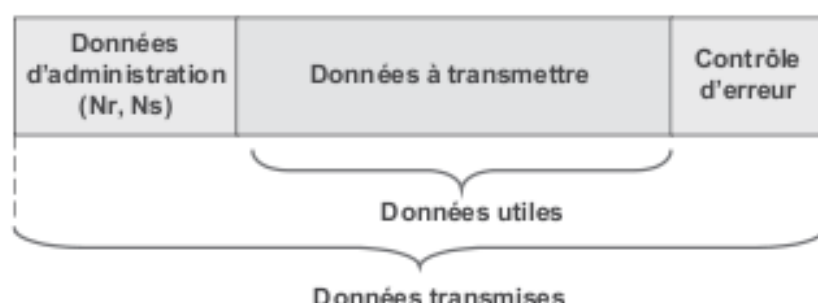


Figure 8.7 La structure de base de la trame HDLC.

■ Les protocoles à anticipation

Les protocoles simples travaillent en mode *Send and Wait*, c'est-à-dire qu'après l'émission d'un bloc de données, ils attendent l'ACK pour envoyer le suivant. Une amélioration substantielle des performances peut être obtenue en émettant les blocs sans attendre la réception des ACK, ce processus se nomme **anticipation**. La figure 8.8 illustre ce principe.

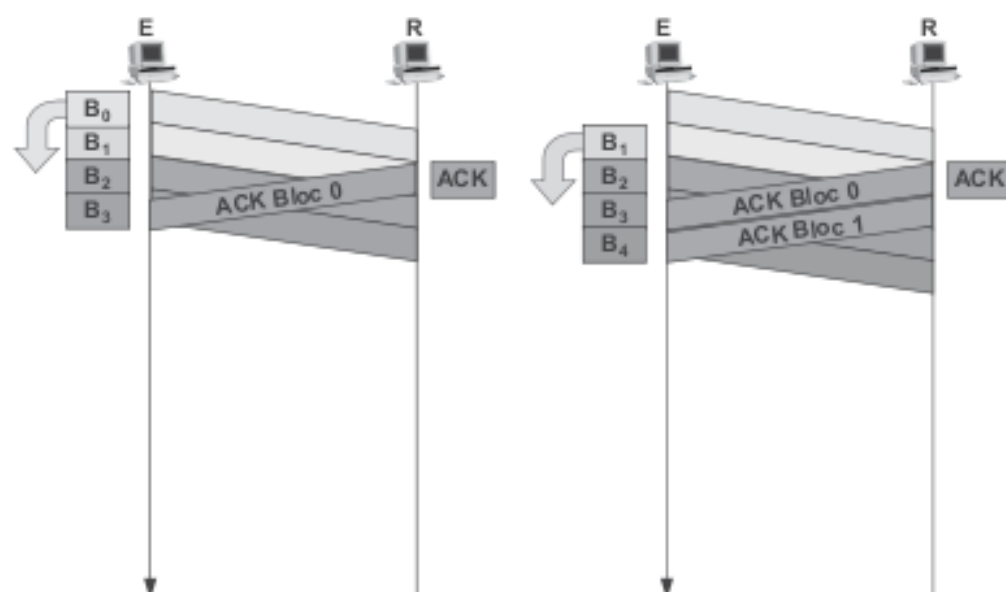


Figure 8.8 Principe des protocoles à anticipation.

L'émetteur émet en continu les blocs de données. Pour autoriser une éventuelle retransmission après erreur (reprise sur erreur), en l'attente de l'acquittement, il mémorise les différents blocs émis (mise en mémoire tampon ou *bufferisation*). À la réception de l'ACK d'un bloc émis, il libère le *buffer*¹ correspondant. La notion d'anticipation est limitée par le nombre de *buffers* que l'émetteur est susceptible de mettre à disposition du protocole.

On appelle fenêtre d'anticipation ou crédit d'émission, notée **W** (*Window*), le nombre de blocs que l'émetteur peut mémoriser en attente d'acquittement. L'efficacité de la transmission est maximale lorsqu'il n'y a pas d'arrêt de l'émission pendant le temps d'attente de l'ACK (émission continue). La taille de la fenêtre optimale correspond donc au nombre de blocs à transmettre pour que l'émission soit continue, la taille optimale de la fenêtre est :

$$W \geq T_a / t_b$$

¹ Le terme *buffer*, en français mémoire tampon, est désormais passé dans le langage courant. Par la suite nous adopterons ce terme, car *bufferisation* est plus évocateur que tamponnage !

Où :

t_b représente le temps d'émission d'un bloc ;

T_a ou temps d'attente représente le temps entre l'émission du premier bit d'un bloc et la réception du dernier bit de son acquittement.

Dans la pratique, T_a est assimilé au temps d'acheminement des données et de leurs acquittements, soit le délai d'acheminement aller et retour y compris le temps de traitement dans la machine cible, ce délai se nomme le **RTT** (*Round Trip Time*).

■ L'acquittement dans les protocoles à anticipation

Chaque bloc n'a pas nécessairement besoin d'être acquitté individuellement. L'acquittement peut être différé et concerner plusieurs blocs. La figure 8.9 illustre ce propos. La fenêtre est de 3, l'acquittement du troisième bloc reçu ($N_r = 3$) acquitte les blocs 0, 1, 2 et autorise l'émission du quatrième bloc qui portera le numéro 3. N_r représente le numéro du prochain bloc attendu. L'acquittement est dit global ou différé.

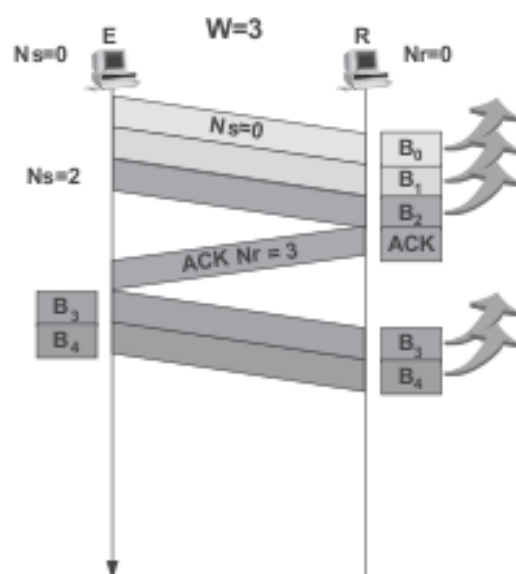


Figure 8.9 Principe de l'acquittement global ou différé.

□ Les protocoles à fenêtre et la politique de reprise sur erreur

Le récepteur délivre aux couches supérieures les blocs reçus au fur et à mesure de leur réception. En cas d'erreur de transmission, deux politiques de reprise sur erreur sont envisageables :

- ▶ le récepteur mémorise les blocs reçus hors séquençement, l'émetteur sur temporisation ou sur demande explicite du récepteur ne retransmet que le bloc erroné (figure 8.10 gauche) ;
- ▶ le récepteur rejette tous les blocs reçus hors séquençement, l'émetteur reprend alors la transmission à partir du bloc perdu, (figure 8.10 droite).

Dans la première hypothèse (figure de gauche), le rejet est qualifié de **rejet sélectif**, la transmission est optimisée mais nécessite des mémoires tampons importantes en réception (*buffers*) et le réordonnancement de tous les blocs. Le nombre de blocs déséquencés pouvant être reçus par le récepteur s'appelle **fenêtre de réception**.

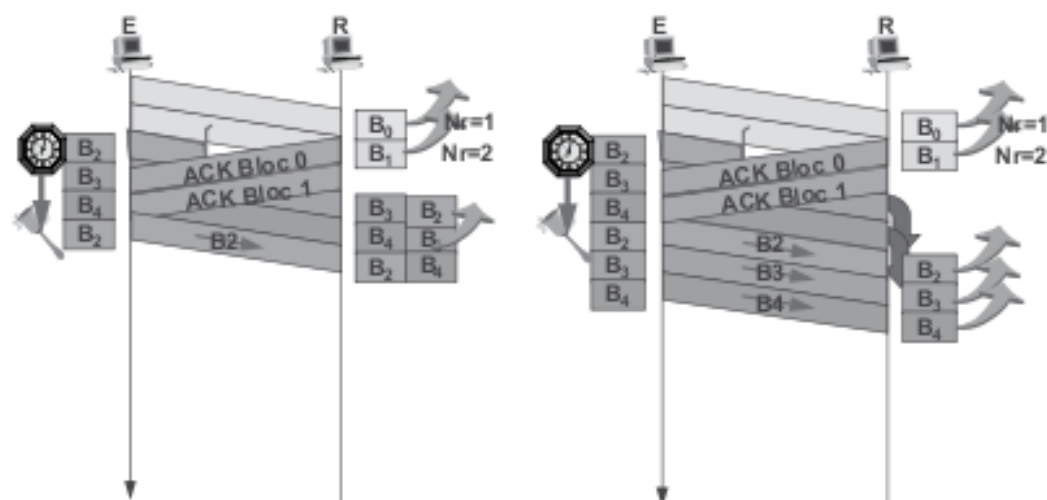


Figure 8.10 Les politiques de reprise sur erreur.

Dans le second cas, la mémoire du récepteur est optimisée, la puissance de calcul du récepteur est minimisée, pas de reséquençement, mais la transmission est pénalisée par la retransmission de tous les blocs. Ce mode de reprise sur erreur est appelé **rejet simple**.

■ Le contrôle de flux

Lors d'une transmission, le destinataire met à la disposition du transfert un certain nombre de mémoires tampons (*buffers*). Le récepteur peut, compte tenu d'autres tâches à réaliser, ne pas libérer ses *buffers* suffisamment rapidement pour en accueillir de nouveaux, les blocs reçus sont alors

perdus. Le contrôle de flux consiste à asservir la cadence d'émission de l'émetteur sur les capacités de réception du récepteur. L'émetteur n'émet alors pas plus de données que le récepteur ne peut en accepter.

La figure 8.11 illustre le principe du contrôle de flux. Dans ce modèle, le récepteur délivre une autorisation explicite à l'émetteur avant l'émission de chaque bloc, le protocole est dit « XON, XOFF ». Le crédit ou fenêtre d'émission est dit de un.

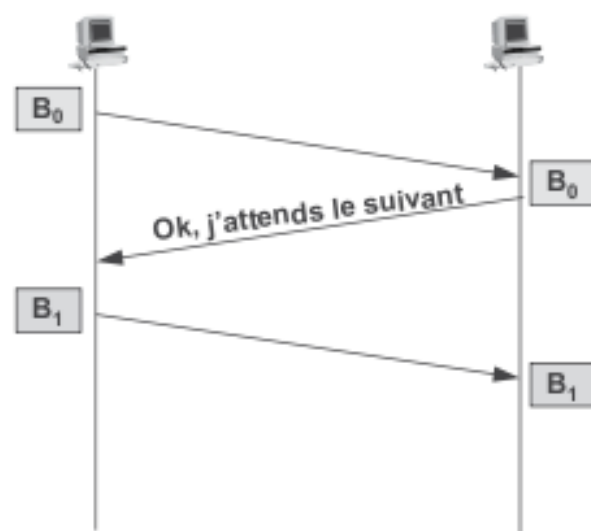


Figure 8.11 Principe du contrôle de flux.

On appelle **crédit d'émission**¹, noté C_t dans la figure 8.12, le nombre de blocs que l'émetteur est autorisé à transmettre. Deux politiques de gestion du contrôle de flux peuvent être envisagées (figure 8.12) :

- ▶ le contrôle de flux est dit **implicite** quand le crédit est prédéterminé (figure 8.11 gauche). Il reste constant durant toute la transmission (fenêtre statique). La transmission est optimisée par rapport au mode *Send and Wait*. Cependant, rien ne permet de garantir que le récepteur pourra toujours recevoir les N blocs du crédit. De plus, la transmission ne bénéficie pas d'éventuelles évolutions des capacités de réception du destinataire. Dans ce mode de fonctionnement, en cas de satura-

¹ La notion de crédit d'émission est souvent confondue avec celle de fenêtre d'anticipation. Bien que les concepts soient proches, la distinction doit être faite. La fenêtre d'émission est un paramètre fixé par l'émetteur alors que le crédit d'émission correspond à une autorisation d'émettre émanant du récepteur.

tion, le récepteur envoie un message de demande d'arrêt des émissions comme dans HDLC ;

- le contrôle de flux est dit **explicite** ou **dynamique** lorsque le récepteur informe en permanence l'émetteur sur ses capacités de réception comme dans TCP (figure 8.12 droite).

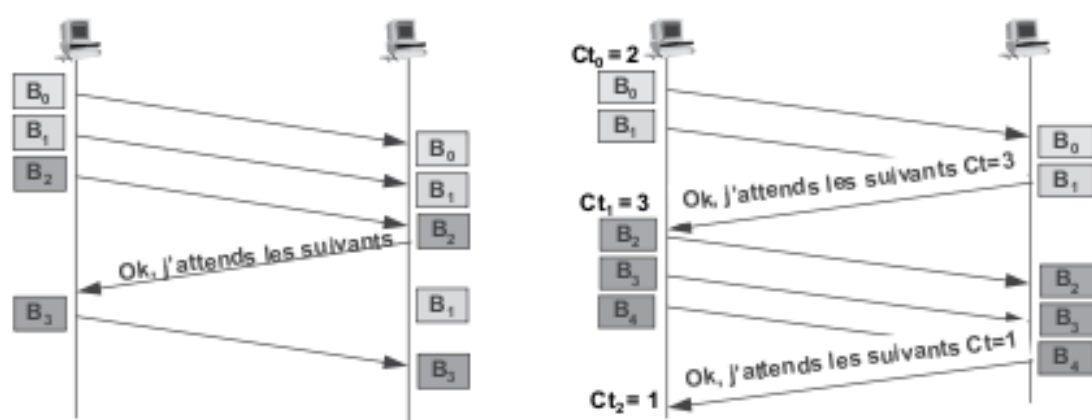


Figure 8.12 Le contrôle de flux par fenêtre.

8.2.2 La structure de la trame HDLC

Le type de la trame émise (information, supervision ou contrôle de liaison) n'est pas distingué par un caractère particulier mais par une combinaison de bits (**protocole orienté bit**). Ce champ de bits est dit champ de commande ou de contrôle. La structure de la trame est donnée par la figure 8.13.

Fanion 01111110	Adresse	Contrôle	Informations	FCS	Fanion 01111110
--------------------	---------	----------	--------------	-----	--------------------

Figure 8.13 La structure générale de la trame.

À l'origine, HDLC était utilisé dans un mode de relation dit « maître/esclave », l'échange ne pouvait avoir lieu qu'entre un terminal (esclave) et la machine maître, seule alors l'adresse de l'esclave était nécessaire. Dans les modes équilibrés (chaque machine peut être maître ou esclave), le champ adresse désigne celui qui est à considérer comme l'esclave dans l'échange.

Le fanion (01111110) délimite la trame : fanion de tête et fanion de queue, ce dernier pouvant faire office de fanion de tête de la trame suivante. La transparence est réalisée selon la technique dite du **bit de bourrage**.

Le champ commande, 8 ou 16 bits selon que les compteurs de trames sont sur 3 ou 7 bits, identifie le type de trame. La figure 8.14 détaille les principales commandes utilisées et précise les combinaisons de bits correspondantes.

Format	Commandes	Réponses	Hex*	Champ Commande			
				8 7 6	5	4 3 2	1
I	INFORMATION		xx	N(r)	P/F	N(s)	0
S	RR		x1	N(r)	P/F	0 0	0 1
	RNR		x5	N(r)	P/F	0 1	0 1
	REJ		x9	N(r)	P/F	1 0	0 1
U	SABM		2F/3F	0 0 1	P	1 1	1 1
	SABME		6F/7F	0 1 1	P	1 1	1 1
	DISC		43/53	0 1 0	P	0 0	1 1
	UA		63/73	0 1 1	F	0 0	1 1
	FRMR		87/97	1 0 0	F	0 1	1 1
	DM		0F/1F	0 0 0	F	1 1	1 1

* les valeurs, exprimées en hexadécimal, dépendent de la position du bit P/F

I	Information	Trame d'information.
RR	Receive Ready	Prêt à recevoir, accusé de réception utilisé lorsque le récepteur n'a pas de trame d'information à envoyer.
RNR	Receive Not Ready	Non prêt à recevoir, le récepteur demande à l'émetteur d'arrêter ses émissions, et acquitte les trames acceptées N(r) - 1.
REJ	Reject	Rejet, demande de retransmission à partir de la trame N(r).
DISC	DISConnect	L'un des ETDD prend l'initiative de la rupture de connexion.
SABM	Set Asynchronous Balanced Mode	Commande permettant le passage en mode équilibré, il n'y a pas de notion de primaire et de secondaire. Chaque station peut émettre sans autorisation.
SABME	Set Asynchronous Balanced Mode	Commande identique à la précédente, mais passage en mode étendu (numérotation modulo 128).
UA	Unnumbered Acknowledge	Acquitte une trame non numérotée.
FRMR	Frame Reject	Informe de la réception d'une trame qui n'a pu être acceptée.
DM	Disconnect Mode	Indique que la station est déconnectée.

Figure 8.14 Les principales commandes d'HDLC.

Enfin, le champ FCS (*Frame Check Sequence*), champ de contrôle d'erreur, contient sur 16 bits le reste de la division polynomiale (CRC) du message transmis (Adresse, Commande, Informations) par le polynôme générateur $x^{16} + x^{12} + x^5 + 1$.

8.2.3 Le fonctionnement d'HDLC

■ Etablissement et rupture de la liaison

La liaison étant dans l'état logique déconnecté (figure 8.15), le primaire demande l'établissement d'une liaison par l'envoi de trames non numérotées (U) de type SABM (mode équilibré ou LAP-B, *Link Access Protocol Balanced*) ou SARM (mode maître/esclave ou LAP). Le secondaire, s'il accepte la connexion, répond par la trame non numérotée UA (*Unnumbered Acknowledge*). La liaison est établie, l'échange d'information peut alors commencer.

La liaison est dans l'état logique connecté. Le primaire émet une demande de déconnexion DISC (figure 8.15), le bit P est positionné indifféremment à 1 ou à 0. Le secondaire accuse réception avec UA, la valeur du bit F correspond à celle du bit P de la trame DISC. La liaison est rompue. L'échange de fanions se poursuit pour maintenir la synchronisation tant que la liaison physique n'est pas rompue.

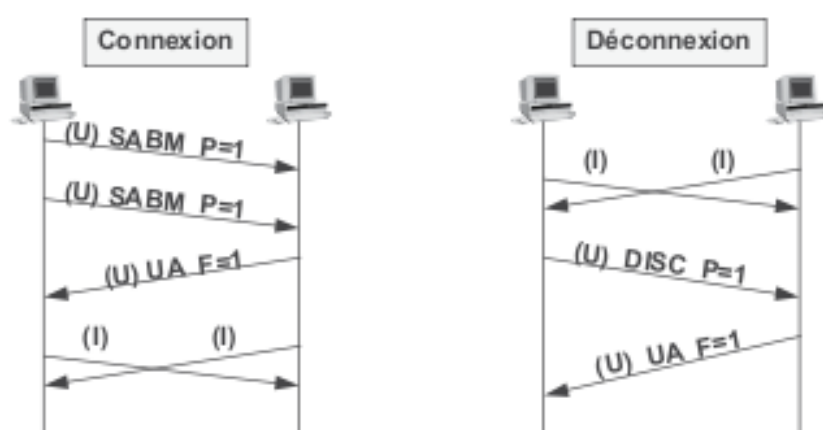


Figure 8.15 La gestion de la connexion sous HDLC.

■ L'échange des données

La figure 8.16 illustre les différentes étapes d'un échange HDLC.

Chaque entité correspondante entretient deux compteurs dits **variables d'état**, le compteur $V(s)$ indique le numéro de la prochaine trame à émettre, le compteur $V(r)$ le numéro de la trame attendue. Après la phase de connexion, les compteurs sont initialisés à zéro de chaque côté. Dans cet exemple, la fenêtre étant de 7, chaque entité a un crédit d'émission de 7 (ligne 1).

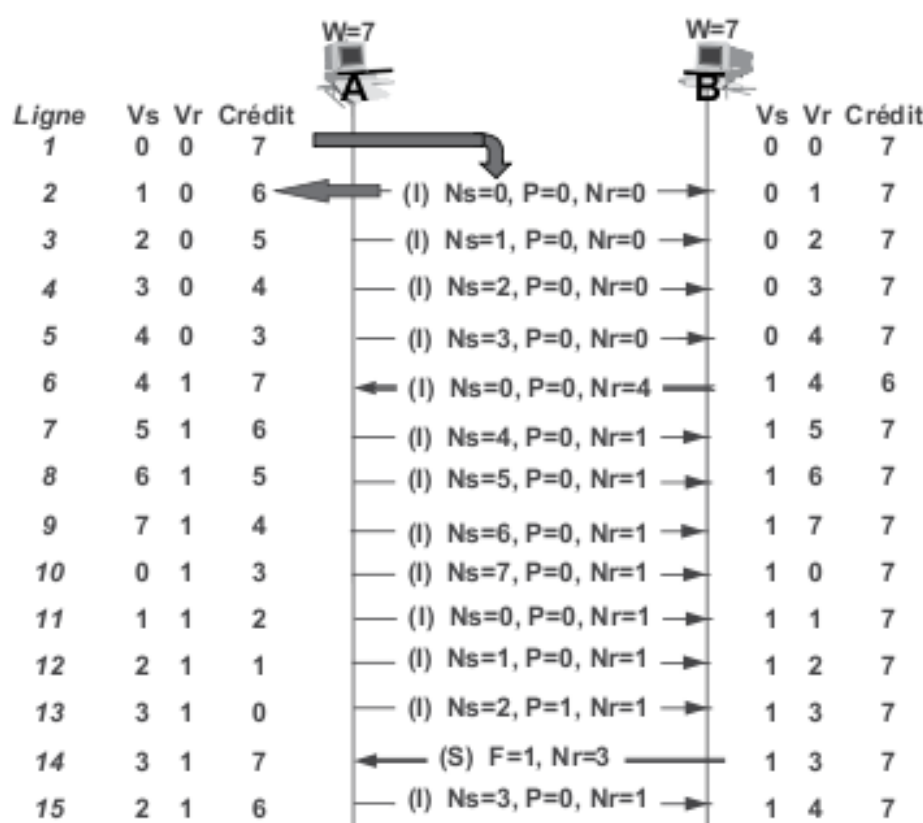


Figure 8.16 L'échange de données et la gestion de la fenêtre.

En ligne 2, la machine A émet une trame, les compteurs $N(s)$ et $N(r)$ contiennent respectivement les valeurs $V(s)$ et $V(r)$ de la ligne 1. Les valeurs $V(s)$, $V(r)$ et crédit de la ligne 2 correspondent aux valeurs, mises à jour après la prise en compte de la trame émise pour la machine A et après réception de celle-ci pour la machine B. C'est-à-dire, que dans la figure, les valeurs des compteurs correspondent toujours aux valeurs mises à jour après réception ou émission d'une trame. Les lignes 3, 4, 5 n'appellent aucun commentaire particulier.

En ligne 6, la machine B émet une trame. Son compteur $V(r)$ contient la valeur de la trame attendue, ici 4, il correspond pour la machine B à un acquittement des $[N(s) - 1]$ trames émises, soit ici les trames 0, 1, 2 et 3. Les mémoires tampons sont libérées, la fenêtre est réinitialisée (crédit de 7).

Dans les échanges « *full duplex* », cette technique d'acquittement simultané à l'envoi de données, dite du *piggybacking*, optimise l'échange de données et évite un blocage de la fenêtre.

L'échange se poursuit, la fenêtre de A s'incrémente. En ligne 12, le crédit n'est plus que d'une trame, il sera nul à l'émission de la trame suivante (ligne 13). La trame émise demande alors un acquittement à B. N'ayant pas de données à envoyer, B acquitte les trames reçues, avec une trame de supervision RR (*Receive Ready*). Il indique à A que cette trame est la réponse à sa demande en positionnant le bit F à 1. La machine A réinitialise alors son compteur de crédit.

■ La gestion des erreurs

La figure 8.17 illustre la reprise sur erreur. Supposons la trame 2 erronée, elle est ignorée par le récepteur. La trame 3 est alors reçue hors séquence, elle est rejetée. La machine B émet une trame de supervision de rejet (REJ, *Reject*) en indiquant à A à partir de quelle trame il doit reprendre la transmission [$N(r) = 2$]. Toutes les trames reçues dont la valeur de $N(s)$ est supérieure à 2 (valeur actuelle de Nr) sont alors rejetées (rejet simple).

La machine A reprend la transmission à partir de la trame 2 ($N(s) = 2$). Si, suite à la trame erronée, A n'avait plus de données à émettre, B n'aurait pas détecté le déséquencelement. C'est A qui, à l'échéance du temporisateur T1, aurait pris l'initiative de retransmettre la trame 2.

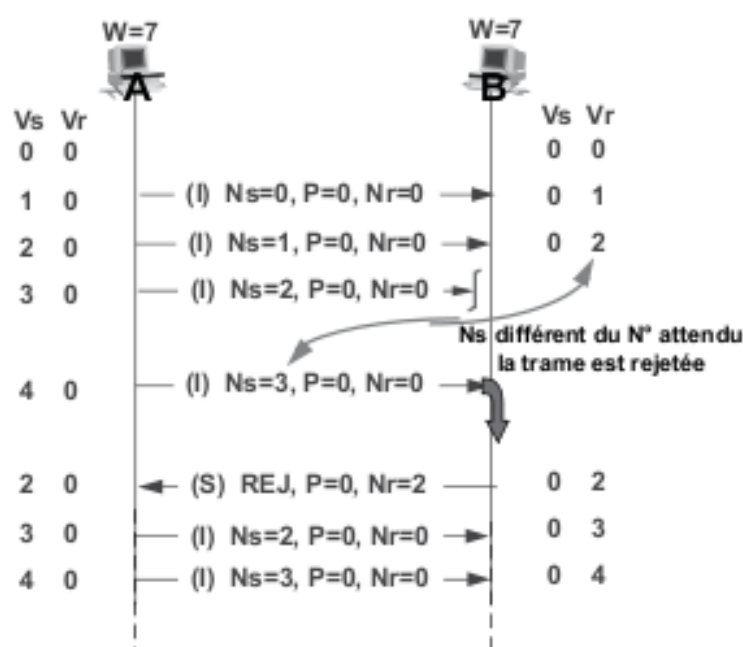


Figure 8.17 La gestion des erreurs.

■ La gestion du contrôle de flux

HDLC utilise le contrôle de flux implicite. La fenêtre est paramétrée à l'installation du logiciel. En cas de saturation des tampons de réception, le récepteur (figure 8.18) rejette la trame en excès et informe A de son incapacité temporaire à accepter de nouvelles données. Il émet la trame « S » **RNR** (*Receive Not Ready*) avec le compteur N(r) positionné au numéro de la trame reçue et rejetée c'est-à-dire la trame à partir de laquelle il faudra reprendre la transmission.

La machine A maintient une activité en émettant régulièrement des trames « S » **RR**, (*Receive Ready*). Lorsque B peut reprendre la réception, il accuse réception à l'aide de la trame « S » **RR**, le compteur N(r) indique la trame à partir de laquelle la transmission doit reprendre.

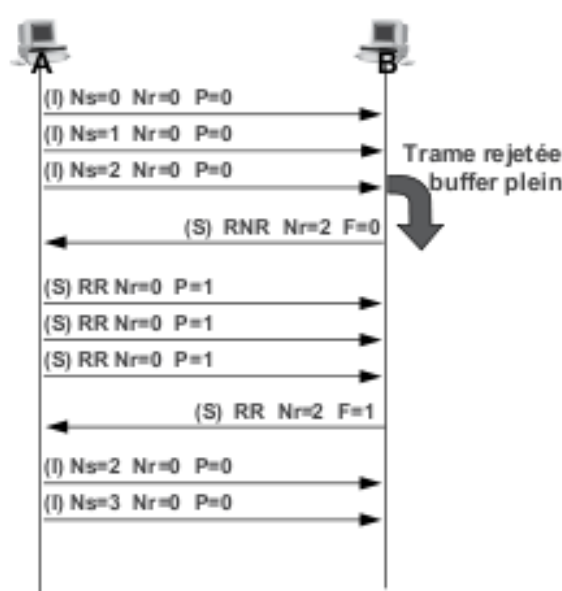


Figure 8.18 La gestion du contrôle de flux.

8.2.4 HDLC et les environnements multiprotocoles

Dans un contexte multi-applications, deux applications d'une même machine sur une même connexion réseau peuvent l'une utiliser un protocole de communication X et l'autre un protocole Y (figure 8.19). Ne distinguant pas les protocoles encapsulés, HDLC ne peut être utilisé que dans un environnement monoprotocole.

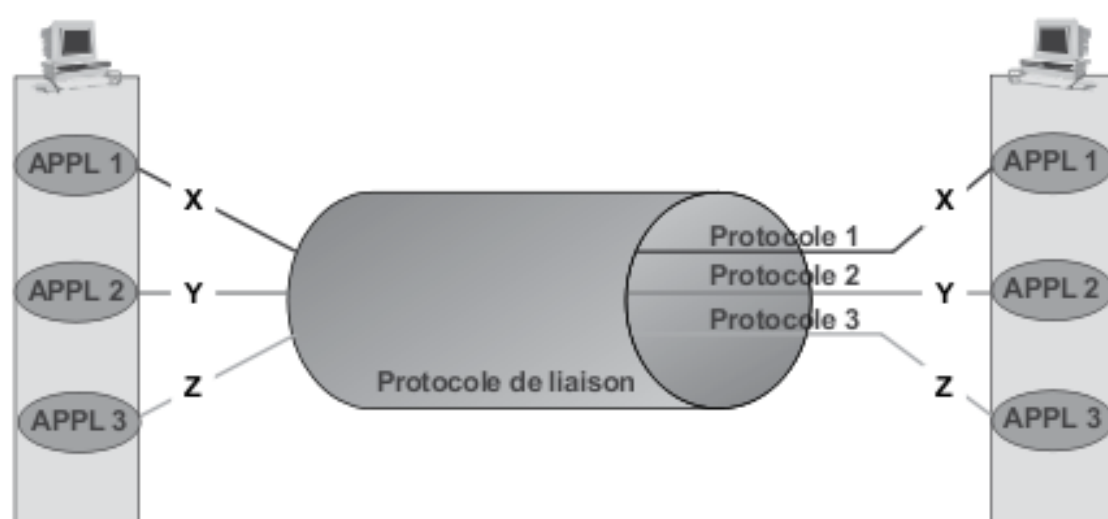


Figure 8.19 La liaison multiprotocole.

Le protocole **PPP** (*Point to Point Protocol*), inspiré de HDLC remédie à cet inconvénient. À cet effet, un champ spécifique : `Protocol_ID` est inséré entre le champ commande et le champ données d'HDLC.

8.3 PPP (Point to Point Protocol)

8.3.1 Généralités

PPP, protocole de liaison point à point (RFC 1548, 1661, 1662), définit un format d'encapsulation dérivé d'HDLC, il comporte un ensemble de protocoles pour (figure 8.20) :

- ▶ le transport de données (PPP proprement dit) ;
- ▶ la négociation des paramètres de liaison (**LCP**, *Link Control Protocol*) ;
- ▶ l'authentification (**PAP**, PPP Authentication Protocol et **CHAP**, Challenge Handshake Authentication Protocol) ;
- ▶ l'obtention des paramètres de configuration de niveau 3 par une famille de protocoles de contrôle de réseau (**NCP**, *Network Control Protocol*).

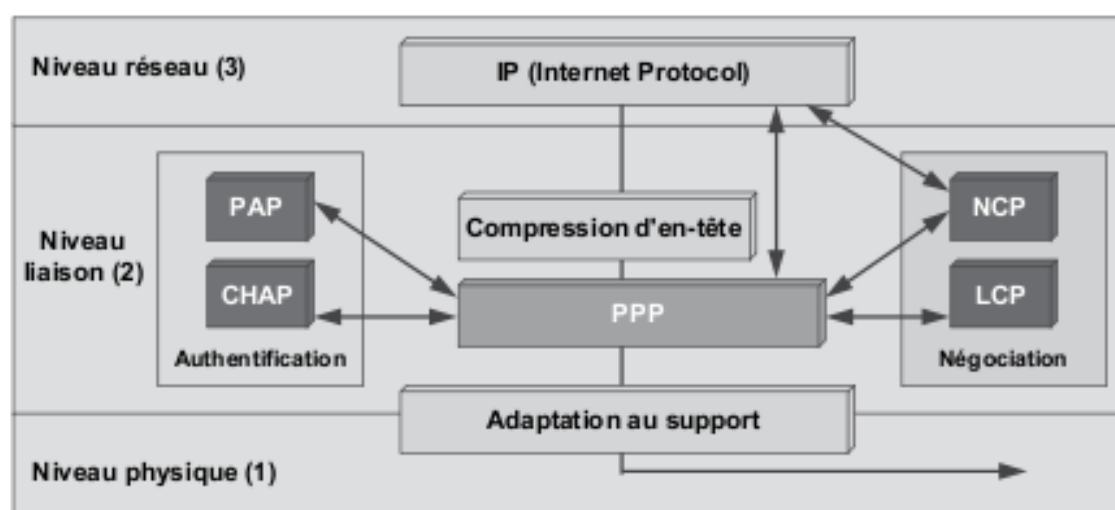


Figure 8.20 PPP et ses sous-protocoles.

8.3.2 L'encapsulation PPP

■ Le format de base

L'encapsulation PPP est une encapsulation à 2 niveaux. Le niveau inférieur, souvent désigné « encapsulation HDLC » est similaire au format de la trame HDLC. Cependant, la présence de tous les champs n'étant pas toujours nécessaire, une négociation préalable entre les entités communicantes autorise un format allégé. L'encapsulation de deuxième niveau permet de distinguer le transport de données applicatives de celles des protocoles complémentaires composant PPP.

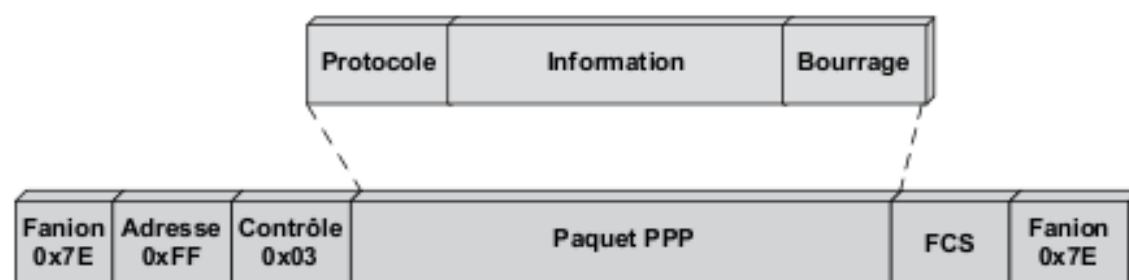


Figure 8.21 Le format générique de la trame PPP.

Hors le champ de données qui encapsule le datagramme IP, les différents champs de la trame PPP sont (figure 8.21) :

- ▶ Le champ **Adresse**, inutile sur une liaison point à point, sa valeur est fixe et vaut 0xFF ;
- ▶ Le champ **Contrôle** a la même signification qu'en HDLC. Si la liaison est fiable, aucun contrôle de séquençement n'est utile (fenêtrage), PPP utilise une trame HDLC « UI » (Trame d'information non numérotée), le champ Commande vaut alors 0x03. Dans les trames UI, les champs adresse et contrôle sont inutiles, ils peuvent être omis (négo-ciation LCP à la connexion).
- ▶ Le champ **Protocole**, sur deux octets, identifie le protocole de niveau supérieur. Lors de la connexion, la longueur de ce champ peut être négociée et réduite à un octet. La figure 8.22 fournit quelques exemples de valeurs du champ Protocole.
- ▶ Enfin, le champ **FCS** (*Frame Check Sequence*) permet la détection d'erreur, son mode de calcul est identique à celui de HDLC.

Valeurs	Protocoles
0x0021	IP
0x002B	IPX
0x002D	TCP/IP (en-têtes compressés, Van Jacobson)
0x002F	TCP/IP (en-têtes non compressés)
0x800F	IPv6
Protocoles complémentaires de PPP	
0xC021	LCP (Link Control Protocol)
0xC023	PAP (Password Authentication Protocol)
0xC223	CHAP (Challenge Handshake Protocol)

Figure 8.22 Les exemples de valeurs du champ Protocole de PPP.

L'une des particularités de PPP est d'avoir été, dès l'origine, imaginé pour fonctionner aussi bien sur une liaison asynchrone¹ que sur une liaison synchrone.

Lorsque le protocole est utilisé sur une liaison synchrone, la transparence au fanion est assurée de manière similaire à HDLC (insertion d'un bit à 0 tous les 5 bits à 1, technique dite du *bit stuffing*). S'il est utilisé sur

¹ Les micro-ordinateurs ne sont équipés d'origine que d'une interface asynchrone.

une liaison asynchrone, la transparence au fanion (0x7E) est réalisée de manière similaire à celle du protocole SLIP. C'est-à-dire (figure 8.23) que le fanion, présent dans le champ de données, est remplacé par la séquence : ESC_PPP, ESC_FLAG soit les valeurs 0x7D, 0x5E. De même, le caractère d'échappement (ESC_PPP, 0x7D) est remplacé par la séquence : ESC_PPP, ESC_ESC soit 0x7D, 0x5D.

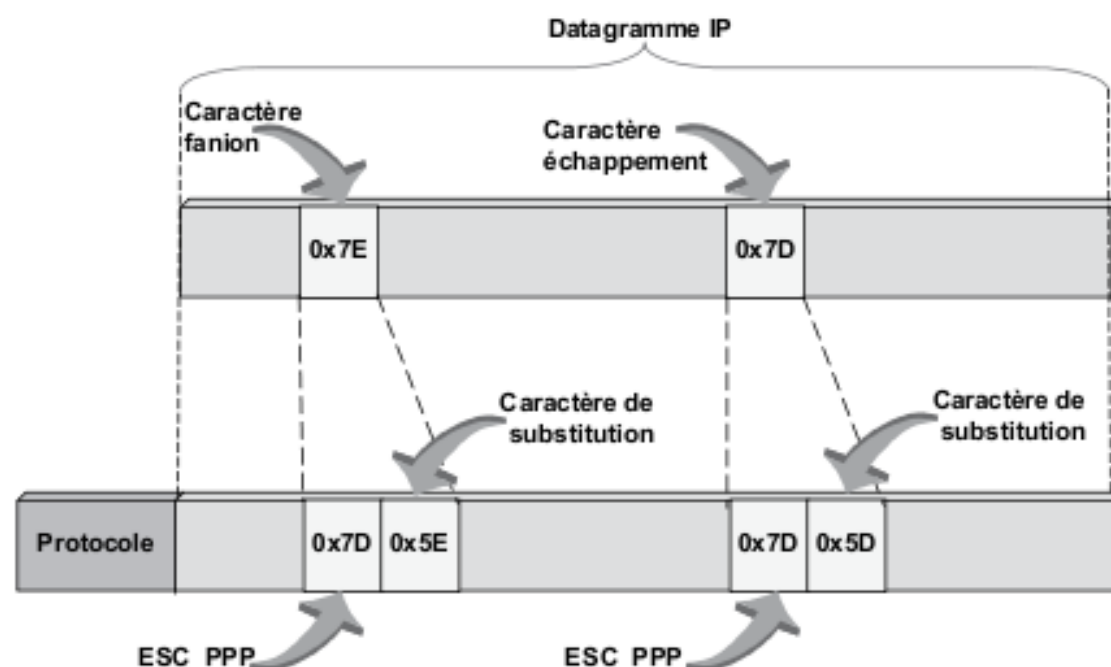


Figure 8.23 Le mécanisme de la transparence dans PPP.

Certains modems utilisent les 32 premiers caractères du code ASCII comme caractères de commande. Pour éviter que ces caractères, lorsqu'ils sont présents dans le champ de données, ne soient considérés par les modems comme des commandes, ils seront, à l'instar du mécanisme de transparence aux fanions, remplacés par des séquences de caractères spécifiques. Les caractères dont la transparence doit être assurée sont indiqués à la connexion lors de l'exécution du protocole LCP (table ACCM, *Asynchronous Control Character Map*).

Par défaut, en mode asynchrone, la table ACCM vaut 0xFFFFFFFF (tous les caractères de commande dans le champ données sont protégés) et en mode synchrone 0x00000000 (aucun caractère n'est protégé).

8.3.3 LCP, Link Control Protocol

Lors de l'initialisation d'un transfert, chaque extrémité de la connexion entreprend une procédure de négociation des paramètres de l'échange par l'intermédiaire du protocole **LCP** (*Link Control Protocol*). La figure 8.24 illustre le format de la trame LCP transportée dans la trame PPP.

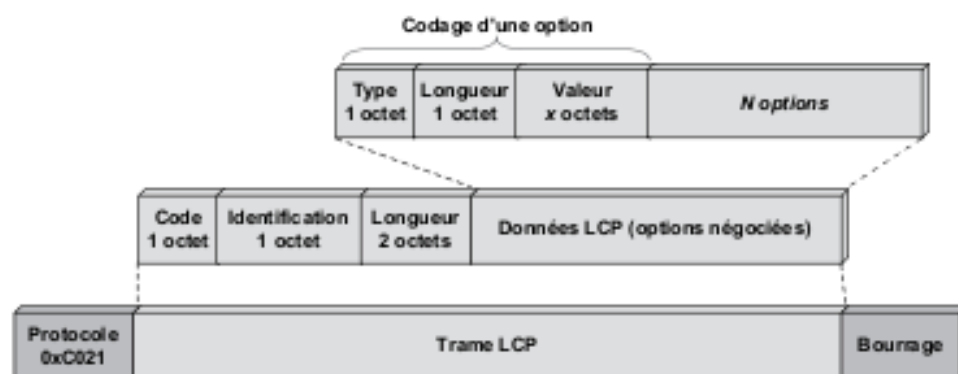


Figure 8.24 La trame LCP.

Le protocole LCP est identifié dans la trame PPP par le code **Protocole_ID** = « 0xC021 » (figure 8.24), le champ **Code** précise le type de trame, le champ **Identification** permet d'associer une requête à une réponse tandis que le champ **Longueur** permet de distinguer les données utiles d'éventuelles données de bourrage. Les options LCP sont codées : **type** d'option, **longueur** du champ valeur et **valeur** des données de l'option. La MTU (code option 1) est, parmi les options négociées, la plus importante.

La MTU (*Maximum Transfer Unit*) définit la capacité d'emport du niveau 2. Pour tenir compte de la taille du champ Protocole, il est nécessaire de définir une valeur de charge utile vue du niveau réseau : la **MRU** ou *Maximum Receipt Unit*, taille maximale du segment admis par le récepteur, valeur par défaut 1 500 octets. La figure 8.25 illustre la relation entre MTU et MRU.

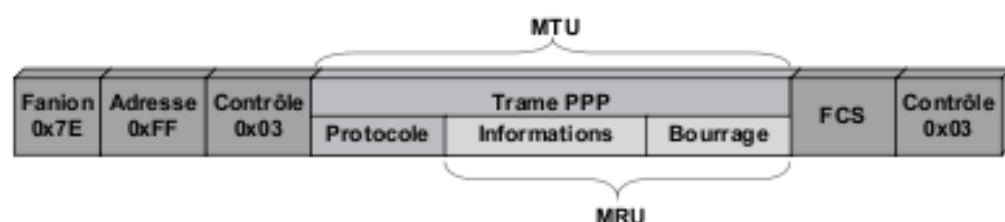


Figure 8.25 La relation entre MTU et MRU.

8.3.4 La sécurisation des échanges

Le protocole PPP est utilisé en mode point à point sur des liaisons permanentes, mais aussi sur des liens temporaires comme, par exemple, une connexion à Internet *via* le réseau téléphonique. De ce fait, l'entité appelée doit pouvoir identifier et authentifier l'appelant avant d'accepter une connexion, c'est le rôle des deux protocoles de sécurisation de PPP : PAP et CHAP.

■ Le protocole PAP, PPP Authentication Protocol

Le protocole PAP échange en clair sur le réseau l'identifiant et le mot de passe de l'appelant. L'échange n'a pour objet que de valider la connexion, ce n'est pas une phase dite de « login » sur un serveur ou une application distante. La figure 8.26 illustre cet échange.

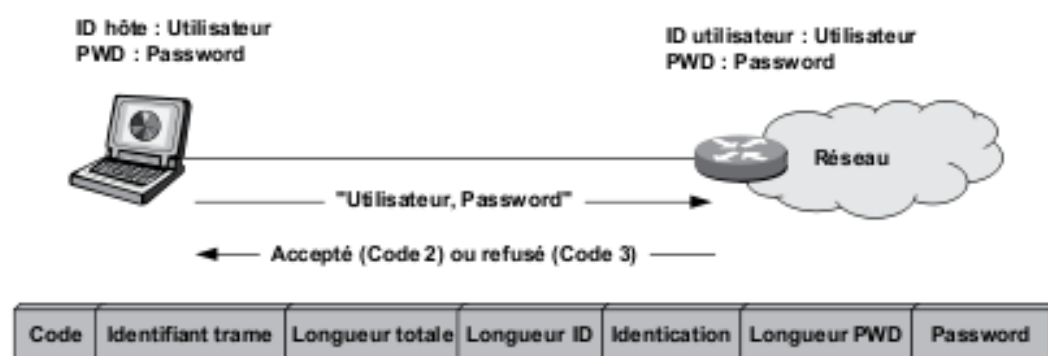


Figure 8.26 L'identification sous PAP et le format de la trame.

Le champ identifiant la trame permet d'associer une requête à sa réponse. Les champs suivants contiennent l'identifiant et le mot de passe de l'extrémité qui s'authentifie.

■ Le protocole CHAP, Challenge Handshake Authentication Protocol

PAP n'est pas un protocole d'authentification fiable. Les mots de passe ne sont échangés qu'une seule fois en début de session et circulent en clair sur le réseau. CHAP est utilisé dès la connexion pour authentifier celle-ci, puis périodiquement pour s'assurer qu'il n'y a pas eu substitution de correspondant. Le protocole CHAP est un protocole en trois temps

(figure 8.27). Après l'établissement du lien, l'appelé envoie un message de test à l'appelant (*Challenge*). Ce dernier chiffre le contenu du message de test et le renvoie. L'expéditeur compare la réponse avec sa propre valeur (système à clé partagée ou secrète). Si le résultat coïncide, l'authentification est acceptée, dans le cas inverse la liaison est rompue.

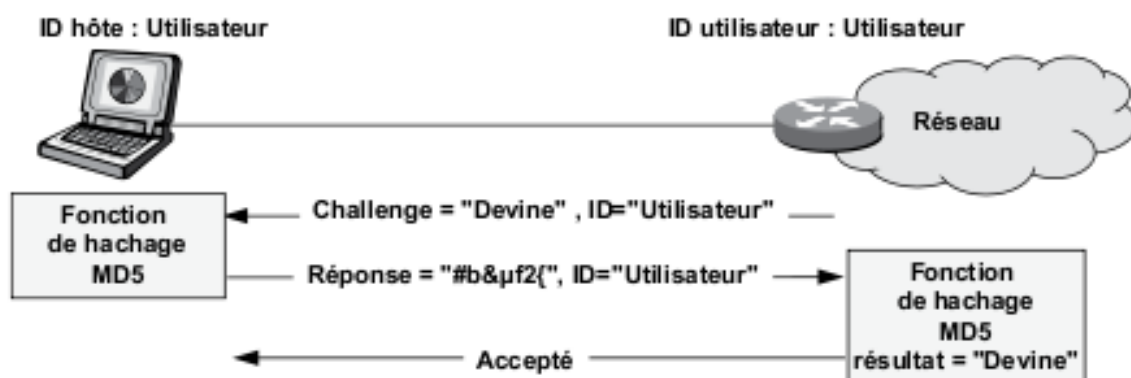


Figure 8.27 Principe de l'authentification en trois temps.

CHAP utilise deux formats de paquets, l'un correspond à l'échange des données du challenge, le second à l'indication du résultat de celui-ci. La figure 8.28 illustre le format de ces paquets.

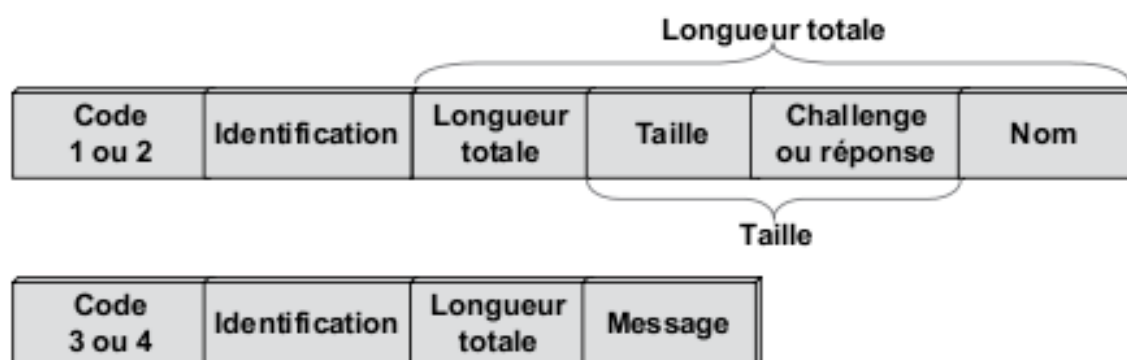


Figure 8.28 Le format des paquets CHAP.

8.3.5 Le protocole NCP, Network Control Protocol

Le protocole NCP est utilisé pour définir les paramètres du niveau réseau. PPP pouvant être interfacé à divers protocoles de niveau 3, NCP est constitué d'un ensemble de protocoles spécifiques à chaque pro-

protocole de niveau 3 notamment IPCP (*IP Configuration Protocol*) pour IP (RFC 1332). Le protocole IPCP est encapsulé dans une trame PPP (protocole=0x8021).

La figure 8.29 illustre le scénario complet d'une connexion PPP. L'entité qui demande l'attribution d'une adresse IP émet sa requête avec le champ Adresse IP à zéro.

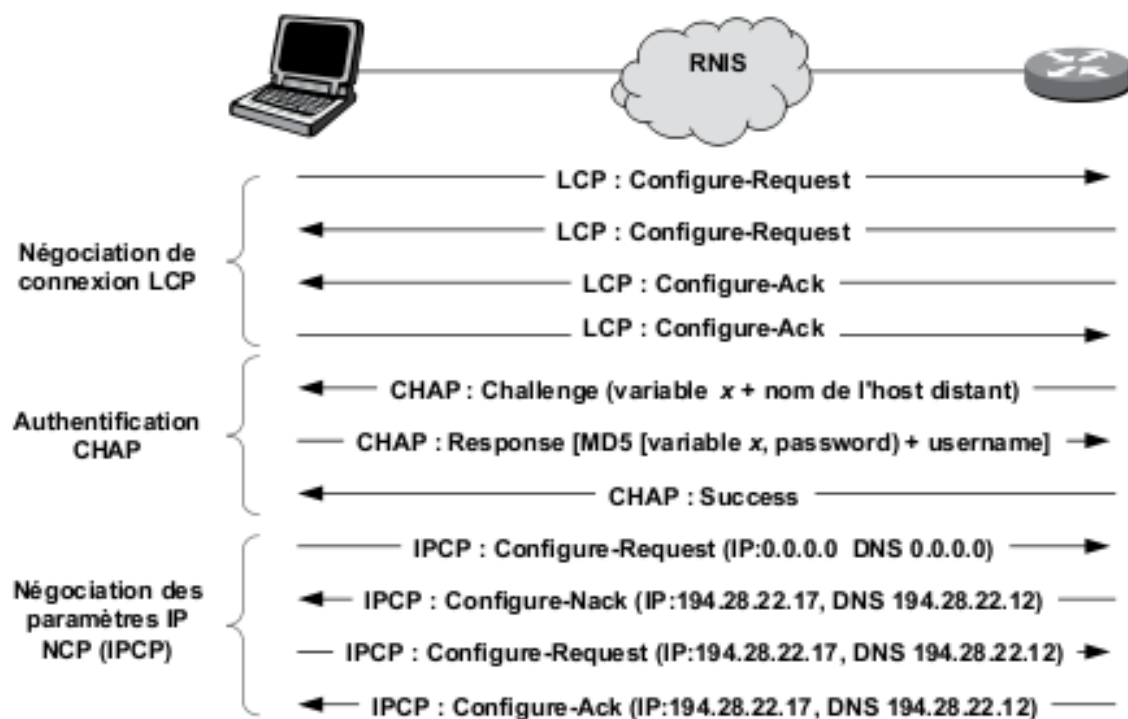


Figure 8.29 Les phases d'établissement d'une connexion PPP.

8.3.6 Les tunnels PPP

Initialement conçu pour être utilisé sur des liens en mode point à point, PPP est rapidement devenu le protocole d'accès à tout type de réseau et notamment au réseau des réseaux (Internet). Avec l'arrivée des connexions à haut débit (x DSL, x Digital Subscriber Line¹), les fournisseurs d'accès ont dû s'équiper d'équipements capables de gérer ces nouvelles connexions.

1 x , indique le type de technique DSL mise en œuvre.

Les **BAS** (*Broadband Access Server*) assurent cette fonction. La connexion entre l'utilisateur et le point d'accès au réseau de l'opérateur (**PoP**, *Point of Presence*) est réalisée, selon le type d'accès, par une adaptation du protocole PPP. L'accès peut alors être réalisé *via* un circuit virtuel ATM (*Asynchronous Transfer Mode*) entre l'abonné et le BAS : PPPoA (PPP over ATM, RFC 2364).

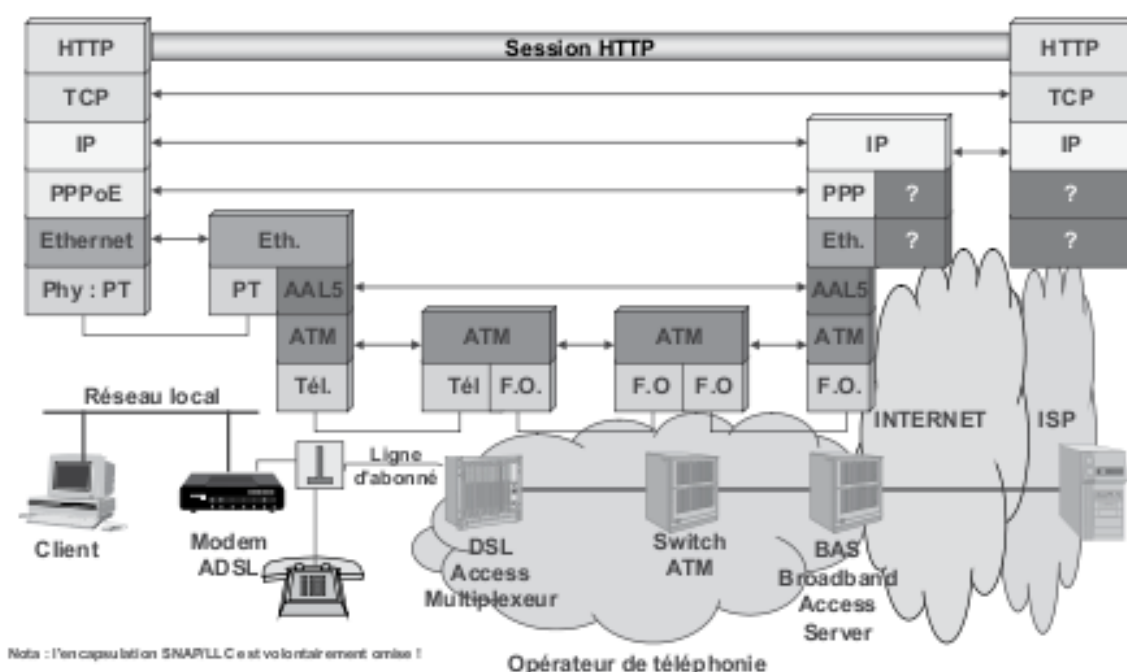


Figure 8.30 Principe d'un accès haut débit à Internet.

Ce mode d'accès a été amélioré en établissant entre le poste client final et le BAS un lien de type Ethernet (PPPoE, PPP over Ethernet, RFC 2516). Dans ce mode de fonctionnement (figure 8.30 où les « ? » symbolisent les protocoles utilisés par les opérateurs), les données IP sont encapsulées dans une trame PPP, elle-même encapsulée dans une trame Ethernet puis dans des cellules ATM (AAL5). Ce mode de fonctionnement simule un pont entre le réseau Internet et le client, il autorise l'utilisation d'une pile protocolaire IP standard et surtout permet de raccorder plusieurs terminaux sur une même connexion.

8.4 Conclusion

HDLC, en version LAP-B, est utilisé dans les réseaux de type X.25, c'est le protocole de liaison par défaut des routeurs Cisco. Les contrôles d'erreur et de flux sont effectués de point à point (nœud à nœud). Cette technique est efficace mais pénalise gravement les performances d'HDLC. L'évolution des techniques réseaux (fibres optiques) rend les supports plus fiables (taux d'erreur plus faible) et autorise une simplification des protocoles. En confiant aux calculateurs d'extrémité (ceux qui sont connectés au réseau), les tâches de contrôle d'erreur et de contrôle de flux, les techniques comme le relais de trames (Frame Relay ou LAP_F) et ATM autorisent des débits de plusieurs dizaines de Mbit/s.

4

Le niveau réseau

9

Le concept de réseau à commutation

9.1 Définitions

Lors de la réalisation d'une liaison de transmission de données, le responsable réseaux et télécoms d'une entreprise recherche la meilleure solution en termes d'efficacité et de coût. Il est envisageable de partager un canal de transmission entre plusieurs utilisateurs par les techniques dites de multiplexage. Cependant, un multiplexeur ne met en relation que des entités prédéfinies, leur connectivité est réduite, la relation est dite de 1 à 1. Les réseaux permettent aussi un partage de la bande passante mais offrent un service de connectivité totale. Un réseau peut être défini comme étant un ensemble de moyens matériels et logiciels géographiquement dispersés destinés à offrir un service, comme le réseau téléphonique, ou à assurer le transport de données. Les techniques à mettre en œuvre diffèrent en fonction des finalités du réseau et de la qualité de service désirée.

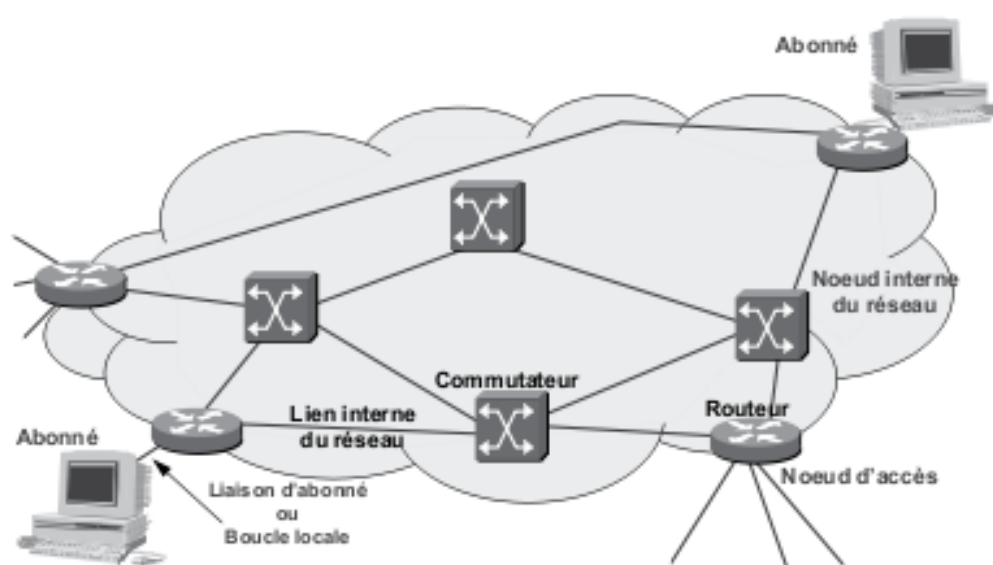


Figure 9.1 Le réseau : ensemble de ressources mises en commun.

Le réseau (WAN) illustré par la figure 9.1 est composé de nœuds (éléments actifs). Les nœuds d'accès, situés à la périphérie du réseau, permettent le raccordement des usagers via une liaison dénommée **liaison d'abonné**. L'ensemble des moyens mis en œuvre pour raccorder un usager est souvent désigné par le terme de **boucle locale**¹. Les nœuds sont généralement des routeurs au point d'accès et des commutateurs au cœur du réseau. Selon la manière dont sont reliés les différents nœuds du réseau ou topologie, on distingue (figure 9.2) :

- ▶ les **réseaux étoiles**, toutes les machines sont raccordées à un nœud central, c'est le cas généralement des réseaux téléphoniques privés où les différents postes téléphoniques sont raccordés au PABX (*Private Automatic Branch eXchange*) de l'entreprise ;
- ▶ les **réseaux hiérarchiques**, composés de différentes étoiles raccordées entre-elles, cette topologie correspond à celle mise en œuvre dans les réseaux locaux ;
- ▶ les **réseaux maillés** dans lesquels il existe plusieurs chemins pour atteindre un point du réseau. Cette topologie est la base de la réalisation des réseaux de type WAN.

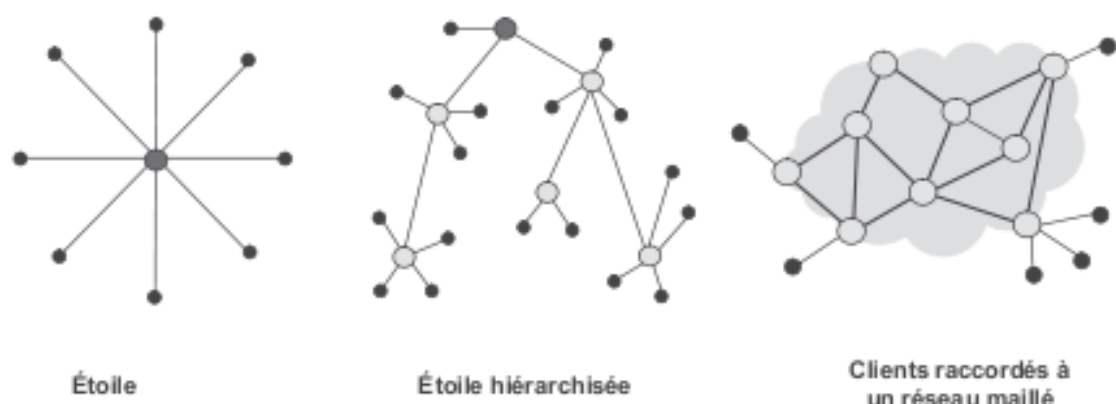


Figure 9.2 Les différentes topologies réseaux.

¹ Une autre définition, plus restrictive, limite la boucle locale à la liaison cuivre qui relie l'abonné au PoP (*Point of Presence*).

9.2 Les réseaux à commutation

9.2.1 Introduction à la commutation

Le concept de réseau à commutation est né de la nécessité de mettre en relation un utilisateur avec n'importe quel autre utilisateur (relation de 1 à 1 parmi n ou interconnexion totale) et de l'impossibilité de créer autant de liaisons point à point qu'il y a de paires potentielles de communicants. En effet, on montre que le nombre total de liens nécessaires, en mode point à point, dans un système comprenant N clients serait de :

$$\text{Nombre de liens} = \frac{N(N-1)}{2}$$

Cette formule montre, s'il en était besoin, la nécessité de trouver un système qui permette, à partir d'une simple ligne de raccordement (liaison d'abonné), d'atteindre simplement tout autre abonné du réseau par simple commutation d'un circuit vers cet autre abonné. Ce système porte le nom de **réseau à commutation**, dans le réseau illustré par la figure 9.3, le commutateur met en relation les utilisateurs A et B.



Figure 9.3 Principe d'un réseau à commutation.

Dans ce contexte où la ressource est rare vis-à-vis de la demande potentielle (si simultanément tous les abonnés du réseau désiraient joindre un autre abonné...), il est indispensable de rechercher des techniques particulières pour optimiser le partage des ressources, c'est l'objectif des techniques de commutation. Selon la technique employée pour « relier » deux utilisateurs, on distingue la commutation de circuits, de messages ou de paquets.

9.2.2 La commutation de circuits

Afin de constituer une liaison de bout en bout, dans les réseaux à commutation de circuits, un lien physique entre une source et une destination est établi par juxtaposition de différents supports physiques (figure 9.4). La mise en relation physique est réalisée par les commutateurs avant tout échange de données et est maintenue tant que les entités communicantes ne la libèrent pas expressément.

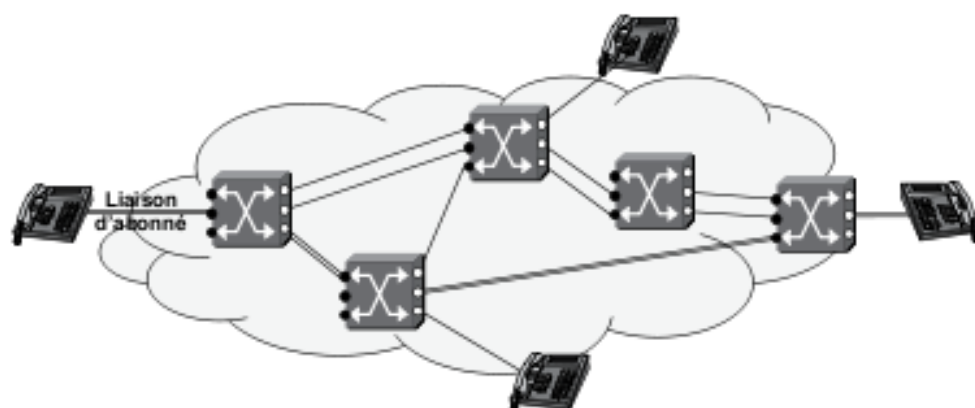


Figure 9.4 Le réseau à commutation de circuits.

La constitution d'un chemin physique, emprunté par la suite par toutes les données transférées, garantit le respect du séquençement des informations.



Figure 9.5 La commutation temporelle.

Archétype des réseaux, la commutation de circuits ou commutation spatiale est aujourd'hui remplacée par une commutation par intervalle de temps (IT) entre des multiplex entrants et des multiplex sortants (commutation temporelle, figure 9.5).

9.2.3 La commutation de messages

En commutation de circuits, la régulation de trafic est réalisée à la connexion ; lorsqu'il n'y a plus de ressource disponible, la connexion est refusée. Contrairement à la commutation de circuits, la commutation de messages n'établit aucun lien physique entre les deux systèmes d'extrémité. Si le lien inter-nœud est occupé, le message est mis en attente (figure 9.6). Chaque bloc d'information (message) constitue une unité de transfert (fichier, écran de terminal...) acheminée individuellement par le réseau. Le message est mémorisé par chaque nœud et retransmis au nœud suivant dès qu'un lien se libère. Le transfert réalisé, le lien est libéré. Assurant une meilleure utilisation des lignes, la commutation de messages autorise un dimensionnement des réseaux inférieur à celui des réseaux à commutation de circuits. En cas de fort trafic, il n'y a pas de blocage du réseau mais seulement un ralentissement (attente de la libération d'un lien). La mémorisation intermédiaire de l'intégralité des messages nécessite des mémoires de masse importantes et augmente le temps de transfert. Les réseaux à commutation de messages ne sont pas adaptés aux applications interactives.



Figure 9.6 Principe de la commutation de messages.

La commutation de messages ne permet qu'un échange **simplex et asynchrone**, c'est plus un service qu'une technique réseau. La commutation de messages est aujourd'hui le support logique des systèmes de messagerie modernes.

9.2.4 La commutation de paquets

La commutation de messages, en monopolisant un circuit durant toute la transmission d'un message, provoque une attente non négligeable

des messages d'autres sources. En découpant le message en fragments (paquets), la commutation de paquets permet l'entrelacement des unités de données (figure 9.7) offrant ainsi, à chaque source, l'impression de disposer en permanence de la voie de communication (multiplexage).

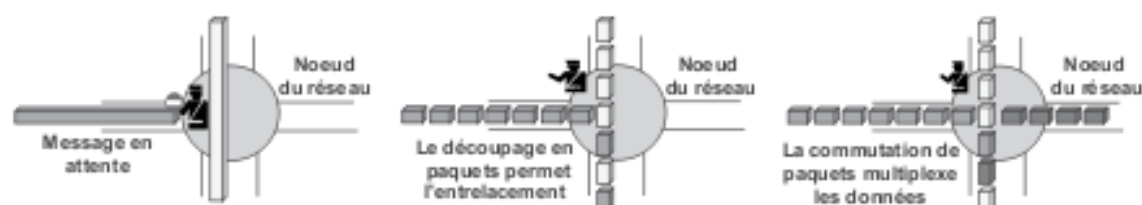


Figure 9.7 De la commutation de messages à la commutation de paquets.

Afin d'optimiser au maximum l'utilisation des liens, chaque paquet est acheminé dans le réseau indépendamment du précédent. Contrairement à la commutation de messages, il n'y a pas de stockage d'information dans les nœuds intermédiaires. Chaque nœud, recevant un paquet, le réémet immédiatement sur la voie optimale au moment de la prise de décision d'acheminement. De ce fait, le séquençement des informations n'est plus garanti. Pour reconstituer le message initial, le destinataire devra, éventuellement, réordonner les différents paquets avant d'effectuer le réassemblage (figure 9.8).



Figure 9.8 Principe de la commutation de paquets.

Ce mode de transfert optimise l'utilisation des ressources, les paquets de différentes sources sont multiplexés sur un même circuit. Cependant, chaque paquet doit contenir les informations nécessaires à son acheminement (adresse ou label). La ressource offerte est banalisée et non attribuée à une communication particulière comme dans la commutation de circuits (figure 9.9).

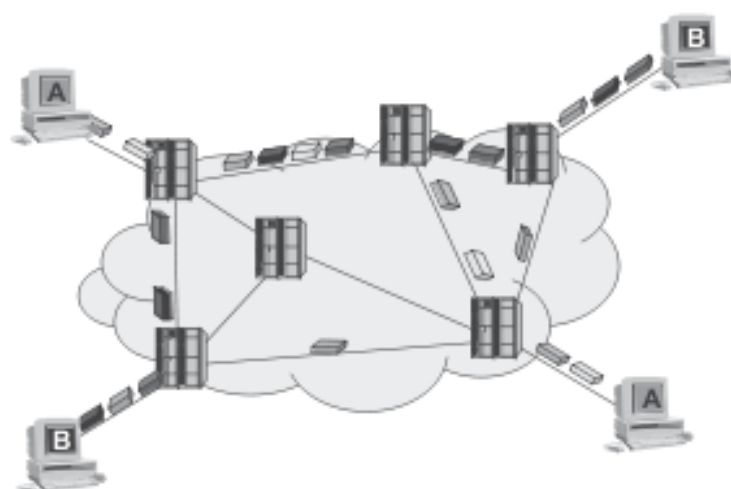


Figure 9.9 Le multiplexage des sources dans le réseau.

Les techniques en relation avec les réseaux en mode paquet sont décrites ci-après.

9.3 Performances des réseaux à commutation

Supposons que dans le réseau à commutation de paquets illustré par la figure 9.10, tous les paquets d'un même message empruntent la même route. En admettant que le temps de transfert sur le support et que le temps de traitement dans les nœuds soient nuls, seul alors le temps d'émission des paquets sur le support intervient pour déterminer les performances.

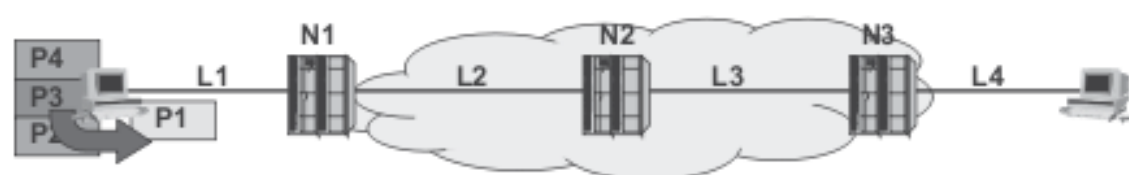


Figure 9.10 La performance d'un réseau à commutation de paquets.

Supposons qu'un message de longueur L (en bits) soit découpé en p paquets et émis sur les différents supports à un même débit de D bit/s. On montre que le temps de transfert dans ce réseau est de :

$$T_p = \left(\frac{L + pH}{D} \right) \left(1 + \frac{N}{p} \right)$$

où H représente les données protocolaires nécessaires à l'acheminement. Cette formule est généralisable pour tous les types de réseau :

- ▶ si on fait $p=1$ et $N=1$, cette formule fournit le temps de transfert dans un réseau en mode circuit ;
- ▶ si on fait $p=1$ et $N > 1$, c'est la commutation de message ;
- ▶ enfin si p et N sont supérieurs à 1, c'est la commutation de paquets.

10 Les réseaux à commutation de paquets

10.1 Du mode datagramme au mode connecté

10.1.1 Le mode datagramme ou mode non connecté

En commutation de paquets, à la réception d'un paquet, le nœud recherche la route optimale (routage). Dans ces conditions, le séquençement des paquets n'est pas garanti. La reprise sur erreur et le contrôle de flux nécessitant une stabilité de route ne sont, par conséquent, pas réalisables. Le réseau est dit *best effort* (pour le mieux), l'unité de données porte alors le nom de **datagramme**. La figure 10.1 met en évidence l'impossibilité d'une reprise sur erreur dans les réseaux en mode datagramme. En effet, le fait pour un nœud intermédiaire de ne pas avoir reçu un paquet ne signifie pas que ce paquet soit perdu, il peut avoir emprunté un autre chemin. Seule l'entité distante a la capacité de détecter valablement la perte d'un paquet (contrôle de bout en bout).

Le cumul des avantages du mode datagramme ou mode non connecté (**CLNS**, *ConnectionLess Network Service*) qui optimise l'utilisation des ressources mais ne garantit pas l'acheminement des données et la commutation de circuits qui permet la reprise sur erreur, garantit le séquençement et autorise le contrôle de flux conduit à définir un nouveau mode qui, sur un réseau en mode paquet, simule un réseau en mode circuit en « balisant » un chemin dans le réseau émulant ainsi un circuit (CV, circuit virtuel). Ce nouveau mode de fonctionnement des réseaux en mode paquet est dit « orienté connexion » (**CONS**, *Connection Oriented Network Service*).

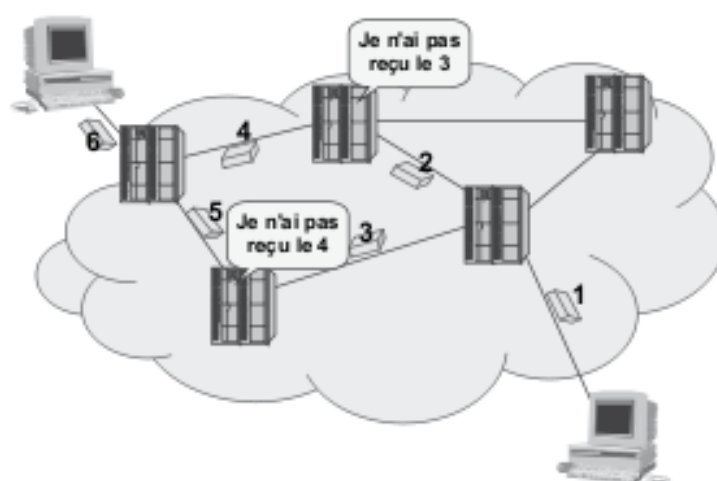


Figure 10.1 La diversité des routes n'autorise pas la reprise sur erreur.

10.1.2 Le mode orienté connexion (CONS)

En commutation de circuits, une liaison physique est préalablement établie avant tout échange de données. En **mode orienté connexion**, une liaison virtuelle est construite par un mécanisme particulier (figure 10.2). Lors de la phase d'établissement de cette liaison, les différentes ressources nécessaires au transfert (*buffers*, voies...) sont réservées. Lorsque l'échange est terminé, une phase de déconnexion libère les ressources. La liaison peut être permanente (**CVP**, Circuit virtuel permanent ou **PVC**, *Permanent Virtual Circuit*) ou établie appel par appel (**CVC**, Circuit virtuel commuté ou **SVC**, *Switched Virtual Circuit*).

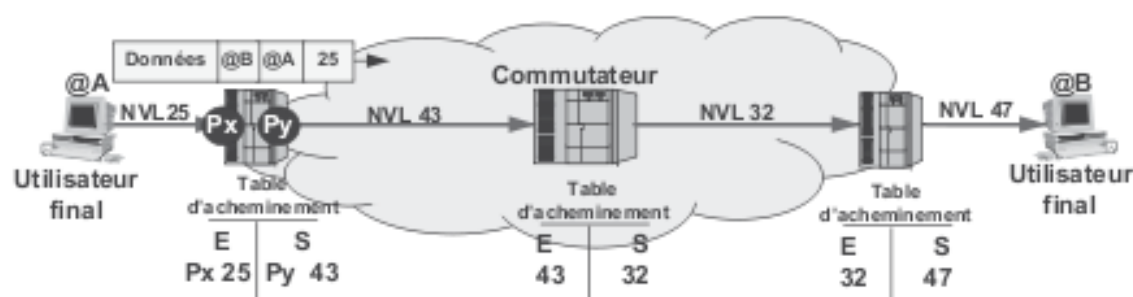


Figure 10.2 Établissement d'un circuit virtuel.

À l'établissement du circuit virtuel, un message spécifique (paquet d'établissement) est acheminé dans le réseau. Chaque nœud traversé associe l'identifiant du paquet et la voie d'arrivée, détermine la voie de sortie et associe celle-ci à un nouvel identifiant (Numéro de voie logique, NVL).

Dans l'exemple illustré par la figure 10.2, la source émet le paquet d'établissement. Celui-ci contient les informations utiles à son acheminement dans le réseau (adresses source et destination) et un label attribué par la source identifie le flux de données. Dans cet exemple, la source attribue un label ou Numéro de voie logique (NVL), ici le numéro 25. Le nœud d'accès au réseau mémorise qu'il a reçu par son port P_x un flux identifié par le NVL 25 ; en fonction de l'adresse de destination et de l'état du réseau, il achemine le paquet sur son port P_y . Compte tenu qu'il avait déjà précédemment identifié sur cette voie 42 autres communications, il substitue au label 25 de la source le label 43 (43^e flux). Il mémorise ces informations dans sa table d'acheminement dite table de commutation. Par la suite, tout paquet entrant par le port P_x et identifié par le NVL 25 sera acheminé sur le port P_y avec le label 43. Chaque nœud jusqu'à destination procède de même. Le circuit virtuel est établi, il résulte de la concaténation des voies logiques identifiées par les labels 25, 43, 32 et 47. À la fin de l'échange, une phase de déconnexion libère les ressources.

10.1.3 Le routage et la commutation

Lorsque la décision d'acheminement est prise en fonction d'une adresse de destination (mode datagramme ou le paquet d'établissement dans le mode connecté), on parle de **routage** ; l'opération est réalisée par un **routeur**. La table d'acheminement est dite **table de routage** (figure 10.3). Cette décision d'acheminement est prise, pour chaque datagramme, par chacun des routeurs traversés (*Hop by hop routing*). Il n'y a aucun contexte d'acheminement mémorisé. De ce fait, ce type de réseau résiste à la défaillance d'un nœud (réseau de type *Soft state* ou réseau sans état). Cependant, la prise de décision par chaque nœud traversé pénalise les performances et donc l'efficacité du transfert de données.

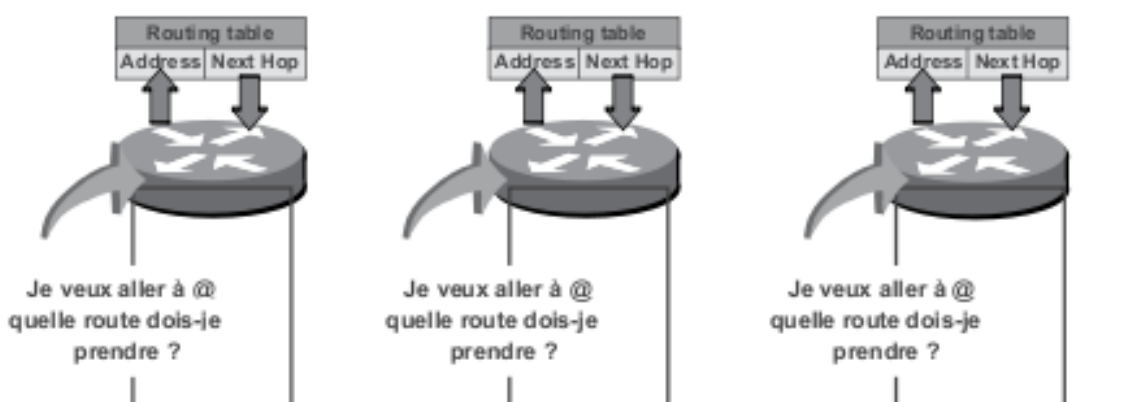


Figure 10.3 Le routage à travers le réseau.

Lorsque l'adresse de destination n'intervient pas dans le processus de décision d'acheminement, on parle alors de commutation. En mode connecté (figure 10.4) préalablement à tout envoi de données, un circuit virtuel est construit par une opération de routage, la **table de commutation** est alors renseignée, les données sont ensuite commutées. Dans la figure 10.4, tout ce qui arrive avec l'étiquette A part du port x et est acheminé sur le port y avec l'étiquette B.

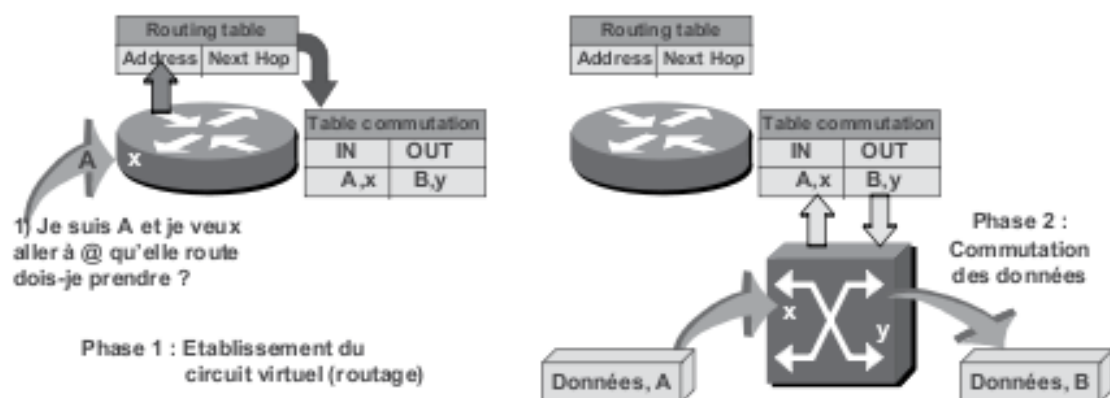


Figure 10.4 Après la phase d'établissement (1), la commutation (2).

Les réseaux en mode datagramme ne disposent d'aucun mécanisme d'établissement de route, ce qui leur confère une grande souplesse et une grande résistance à la défaillance. Cependant, le mode connecté, en garantissant le séquençement des informations et en optimisant le processus d'acheminement, présente des avantages certains, en particulier pour le

transfert des flux multimédia. Aussi, a-t-on imaginé un protocole réseau qui réalise une synthèse entre ces modes en offrant dans un environnement datagramme les performances du mode connecté : c'est l'objectif de **MPLS** (*MultiProtocol Label Switching*) qui implémente un acheminement commuté dans un environnement de type datagramme.

11

Les techniques réseau

11.1 La notion d'adressage

Pour localiser sans ambiguïté un utilisateur final, il faut pouvoir identifier (figure 11.1) :

- ▶ le réseau auquel il est connecté ;
- ▶ le point par lequel il est raccordé au réseau, ce point identifie aussi l'installation locale de l'abonné ;
- ▶ le système cible dans l'installation locale.

L'ensemble de ces informations constitue l'adresse réseau dite aussi adresse logique. Les deux premiers champs de la figure 11.1 permettent de localiser l'installation de l'abonné, ils constituent l'adresse réseau proprement dite du destinataire, la structure en est généralement de type hiérarchique. Le troisième champ identifie le destinataire dans l'installation finale, c'est un simple identifiant sans signification particulière, cet adressage local est dit **à plat**.

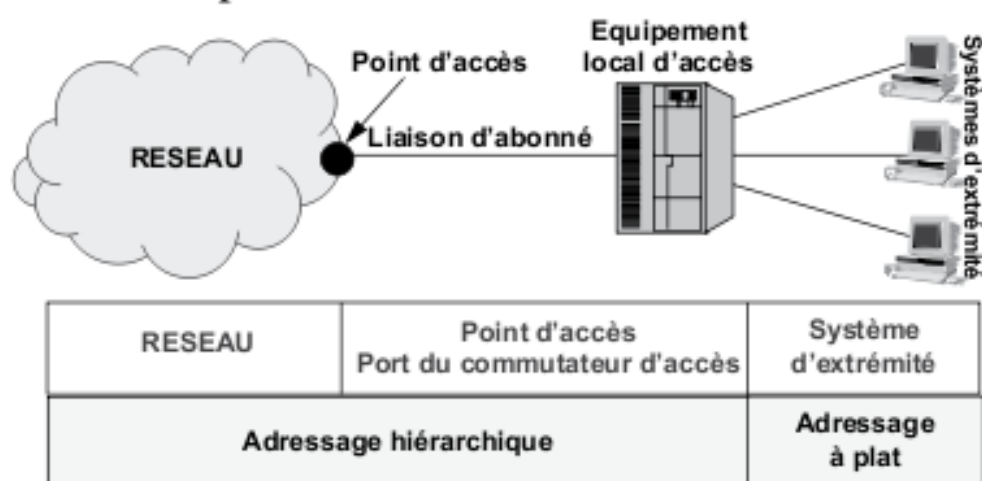


Figure 11.1 Les composants d'une adresse.

Cet identifiant ou adresse logique permet de joindre une machine en tout point du réseau, complément de celle-ci ou parfois partie de celle-ci, l'adresse physique identifie le système d'extrémité ou l'interface de raccordement de la machine cible. Ces notions sont complétées par un mnémonique ou adresse symbolique qui permet de désigner par un nom, plus facilement mémorisable, un processus ou une machine cible.

11.1.1 L'adressage à plat ou global

Dans ce type d'adressage, l'adresse correspond à un numéro unique attribué sans aucune règle de structuration. Cet adressage est, par exemple, celui utilisé dans les réseaux locaux. À chaque entité raccordée est attribué un numéro différent et sans relation avec n'importe quel autre numéro (adresse) du réseau. D'origine Xerox et normalisé par l'IEEE¹ (figure 11.2), cet adressage destiné à distinguer les différents nœuds d'un même segment de réseau est souvent désigné sous le terme d'adressage **MAC** (*Medium Access Control*) ou encore d'adresse physique.



Figure 11.2 L'adressage MAC ou IEEE (réseaux locaux).

L'adressage MAC comporte deux champs. Le premier, champ attribué par l'IEEE, désigne le constructeur (**OUI**, *Organizationally Unit Identifier*) de l'interface réseau (**NIC**, *Network Interface Card*). Le second champ correspond à un numéro séquentiel attribué par le constructeur qui doit en garantir l'unicité.

L'adresse MAC peut identifier un point de raccordement unique (cas général), elle est alors dite *unicast*. Elle peut aussi désigner un groupe de

¹ Notons que l'IEEE a récemment introduit la notion d'identifiant d'interface sur 64 bits (24+40), cet identifiant d'interface est désigné sous le terme EUI-64 (End-User Identifier).

machines raccordées à un même segment du réseau, elle est alors, dite *multicast*. Une adresse MAC spécifique désigne toutes les machines d'un même réseau physique, cette adresse est dite adresse de diffusion généralisée ou *broadcast* (figure 11.3).

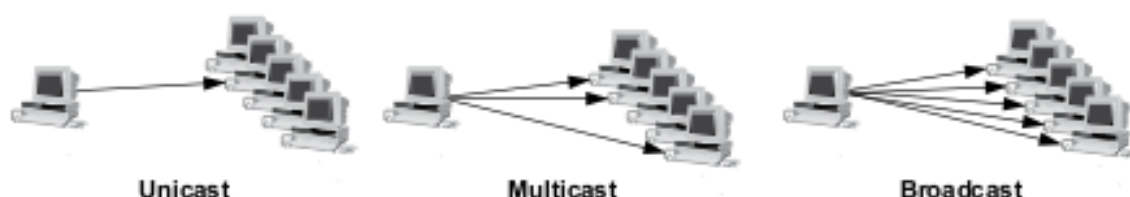


Figure 11.3 L'adressage et les points adressés.

11.1.2 L'adressage hiérarchique

Utilisée dans les grands réseaux d'interconnexion, l'adresse hiérarchique identifie un point d'accès au réseau. Son contenu est significatif, il désigne le réseau et les nœuds de ce réseau participant à l'acheminement des informations. Chaque nœud ne traite que la partie d'adresse correspondant à son niveau.

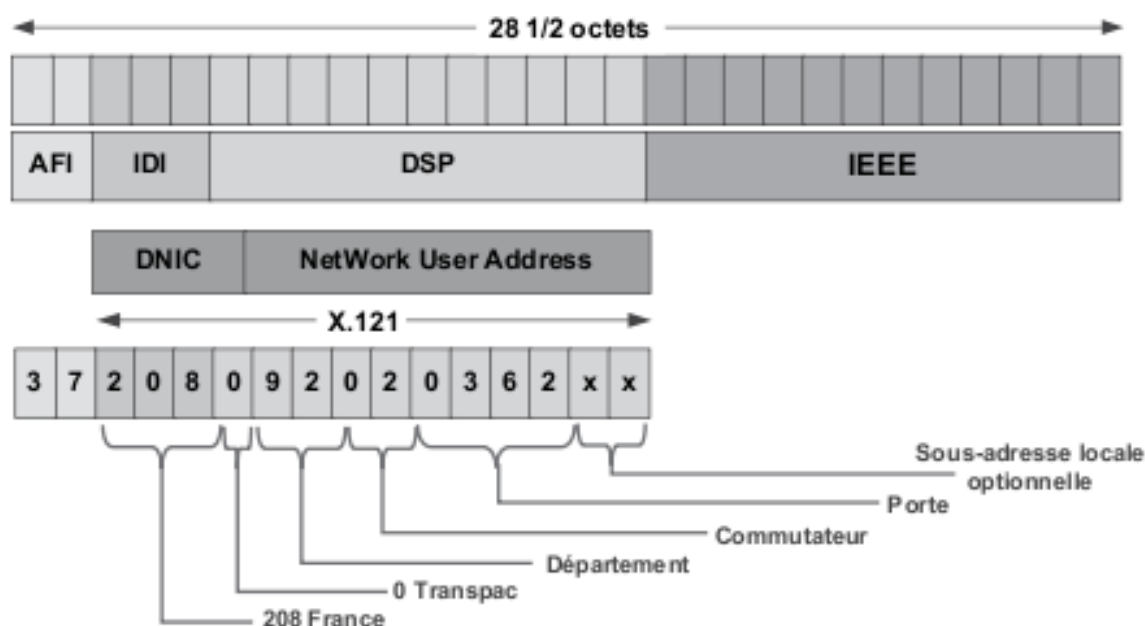


Figure 11.4 L'adressage X.121, application à l'ancien réseau X.25 de Transpac.

Par exemple, l'adressage défini par l'ISO dit adressage **NSAP** (*Network Service Access Point*) représenté figure 11.4 comporte plusieurs champs :

- ▶ L'**AFI** (*Authority Format Identifier*) désigne l'autorité gestionnaire du domaine d'adressage et le format de représentation de l'adresse. La valeur 37 indique que l'adresse qui suit est au format X.121 et est codée en DCB¹ ;
- ▶ L'**IDI** (*Initial Domain Identification*) identifie le domaine d'adressage. Par exemple, dans la norme X.121 (AFI = 37), le numéro 208 est affecté à la France, le 2 représentant l'Europe ;
- ▶ **DSP** (*Domain Specific Part*) correspond à l'adresse effective de l'abonné ;
- ▶ Cette adresse peut éventuellement être complétée par l'adresse du terminal dans l'installation d'abonné, ici nous avons joint à cette adresse l'adresse IEEE du terminal.

11.1.3 Le nommage

La notion de nommage ou adressage symbolique est complémentaire de celle d'adressage, l'un désigne l'objet, l'autre précise sa localisation. Indépendamment du fait qu'il est plus aisé de manipuler des noms que des adresses, l'avantage du nommage est essentiellement de dissocier l'objet de sa localisation géographique. Le déplacement de l'objet nommé est transparent à l'utilisateur. De manière similaire à l'adressage, le nommage utilise deux modes de représentation :

- ▶ Le **nommage à plat ou horizontal**, ce type de nommage impose une démarche rigoureuse pour garantir l'unicité d'un nom sur l'ensemble du réseau. NetBios (Windows), protocole allégé mis en œuvre dans les réseaux locaux, utilise un nommage à plat.
- ▶ Le **nommage hiérarchique ou arborescent**, plus souple, organise le nommage en domaines. Cette technique autorise une représentation des objets calquée sur l'organisation de l'entreprise. Ce mode de représentation et d'administration convient parfaitement à la gestion d'un annuaire très important comme celui d'Internet (figure 11.5).

¹ DCB, Décimal Codé Binaire, dans cette forme de représentation des données, chaque quartet d'un octet code un chiffre décimal, ce qui permet un codage et décodage facile.

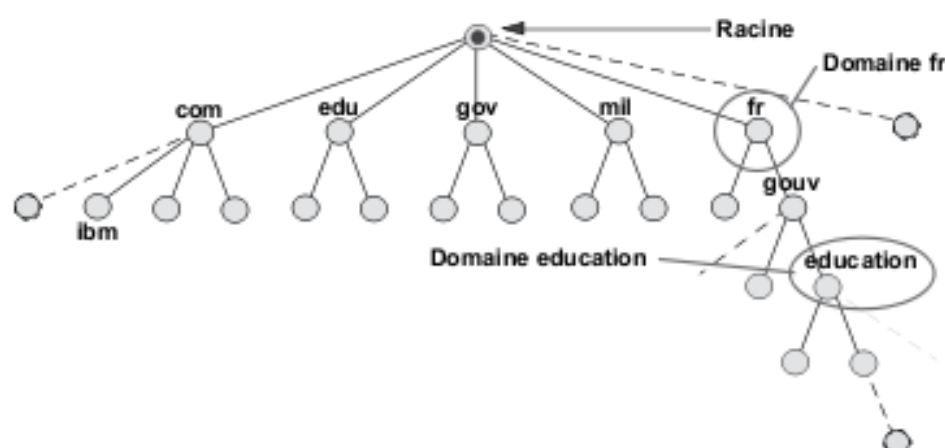


Figure 11.5 L'arbre de nommage d'Internet.

11.2 La segmentation et le réassemblage

11.2.1 La notion de MTU

Lors du transfert d'un bloc de données dans un réseau, chaque élément du réseau (routeur ou commutateur) doit mémoriser les blocs en entrée, les traiter et les délivrer à la file d'attente de sortie. Ces différents traitements nécessitent de la mémoire. La ressource étant limitée, il est nécessaire de fixer une taille maximale aux unités de données admises dans un réseau.

On appelle **MTU** (*Maximum Transfer Unit*) ou unité de transfert maximale, la taille maximale des données admises dans un réseau en-tête compris. Si un bloc a une taille supérieure à la MTU, il devra être fragmenté en plusieurs blocs ou fragments pour pouvoir être acheminé dans le réseau.

11.2.2 La segmentation et le réassemblage

Dans les réseaux en mode non connecté, les fragments sont susceptibles d'arriver sans respect de l'ordonnancement. Le réassemblage ne peut être réalisé dans le réseau, c'est le destinataire qui devra reconstituer l'unité de données d'origine. Pour garantir le réassemblage correct du message initial, il est nécessaire d'identifier tous les fragments d'un même datagramme et de les numéroter pour les réordonner.

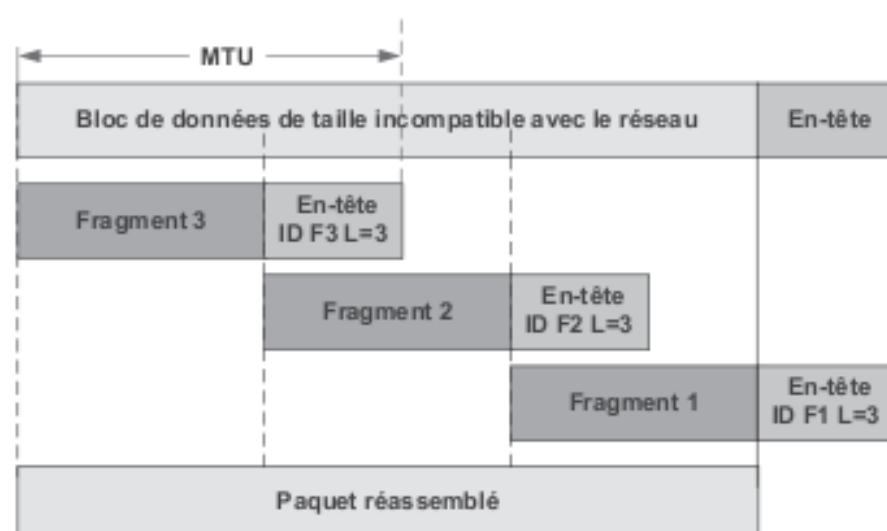


Figure 11.6 Les informations de fragmentation en mode non connecté.

Chaque fragment (figure 11.6) comporte les informations nécessaires à son acheminement (adresse). Une donnée d'identification est recopiée dans chaque fragment (ID) d'un même datagramme. Le réassemblage nécessite aussi de connaître la longueur totale du paquet d'origine (L) et de disposer d'information sur l'ordonnancement (F1, F2...). Dans un environnement datagramme, outre le temps nécessaire aux opérations de fragmentation, la reprise sur erreur dans le réseau étant impossible, la perte d'un seul fragment implique la retransmission de tout le datagramme. Pour ne pas pénaliser le réseau, les protocoles en mode non connecté offrent des services de découverte de la MTU (*Path MTU Discovery*).

11.3 Le contrôle de congestion

11.3.1 Définition

Basé sur un trafic sporadique et aléatoire, le partage statistique des ressources d'un réseau fragilise celui-ci. À une augmentation du trafic soumis, correspond une augmentation du temps d'attente avant traitement dans les nœuds. Vu du destinataire, le débit diminue, le temps de transit dans le réseau croît (congestion légère). Les paquets retardés peuvent, dans ce cas, ne pas être acquittés dans les délais, ce qui provoque leur retransmission et contribue à augmenter la charge du réseau, plus de

paquets ne sont pas acquittés à temps, plus de réémissions inutiles..., les files d'attente débordent..., le réseau s'effondre, c'est la congestion sévère (figure 11.7).

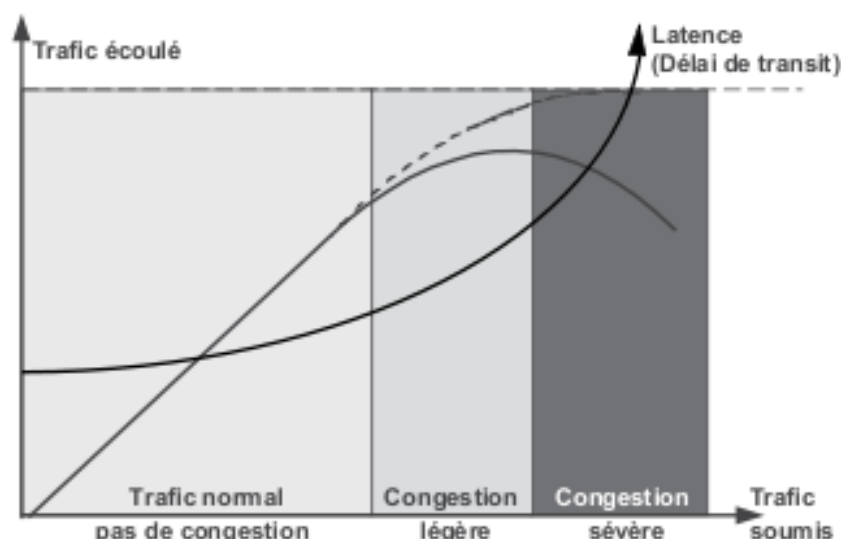


Figure 11.7 L'écoulement du trafic dans un réseau.

En présence d'une surcharge du réseau, les mécanismes de reprise sur erreur des protocoles ont tendance à réagir ensemble. L'indépendance des sources n'est plus vraie, la congestion s'installe.

Indépendamment du blocage du réseau, si on veut garantir une certaine qualité de service, il est nécessaire de mettre en œuvre des mécanismes spécifiques pour d'une part, prévenir l'état de congestion et, d'autre part, si celui-ci apparaît, le résoudre (guérison). Ces mécanismes constituent le contrôle de congestion.

Les notions de contrôle de flux et de contrôle de congestion sont différentes. Le contrôle de flux s'intéresse aux échanges entre deux nœuds alors que le contrôle de congestion cherche à limiter le nombre de paquets en transit dans le réseau (figure 11.8). Cependant, en limitant la longueur des files d'attente dans les nœuds intermédiaires, le contrôle de flux participe à la prévention de la congestion.

En synthèse, le contrôle de flux asservit le débit de la source en fonction des capacités de réception du destinataire (mode point à point ou de bout en bout), alors que le contrôle de congestion limite le débit de la source en

fonction des capacités de transport du réseau. Ainsi, le contrôle de flux est un mécanisme qui consiste à éviter la perte de données par saturation des *buffers* du destinataire, alors que le contrôle de congestion vise à ne pas perdre de données par saturation du réseau.

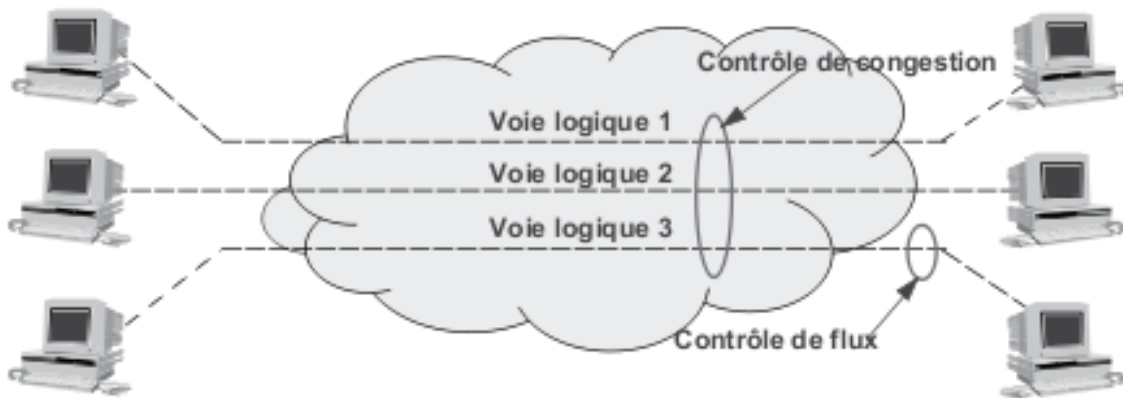


Figure 11.8 Distinction entre contrôle de flux et contrôle de congestion.

11.4 L'acheminement

11.4.1 Définition

Acheminer les informations dans un réseau consiste à assurer le transit des blocs d'un point d'entrée à un point de sortie désigné par son adresse. Chaque nœud du réseau comporte des tables, dites **tables d'acheminement** couramment appelées **tables de routage**, qui indiquent la route à suivre pour atteindre le destinataire. En principe, une table de routage est un triplet <Adresse destination>/<Route à prendre>/<Coût>.

11.4.2 Les protocoles de routage

■ Les différents modes de routage

□ Routage statique ou routage fixe

Le routage statique consiste à construire, dans chaque nœud, une table indiquant, pour chaque destination, l'adresse du nœud suivant. Cette table est construite par l'administrateur du réseau lors de configuration

du réseau et à chaque changement de topologie. Simple, le routage fixe assure, même lorsque le protocole réseau est en mode datagramme, le maintien en séquence des informations. Aucun bouclage de chemin n'est à craindre, mais il n'existe pas de solution de secours en cas de rupture d'un lien. La figure 11.9 illustre le contenu des tables de chacun des nœuds pour joindre b_2 .

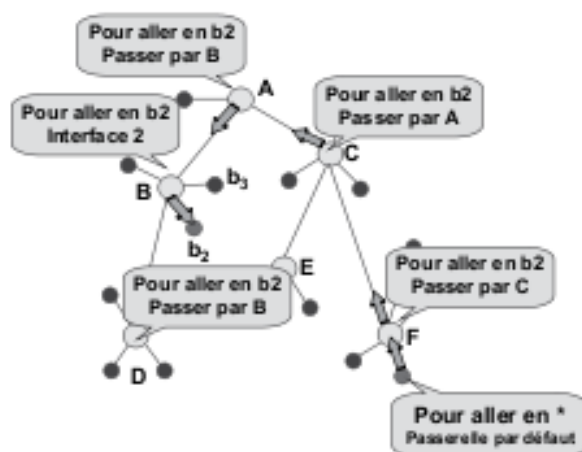


Figure 11.9 Le routage statique.

□ Routage par diffusion (de 1 vers n)

Lorsque l'information doit être routée simultanément vers plusieurs destinataires ou groupe d'utilisateurs, il faut dupliquer le message en autant d'exemplaires que de destinataires. Cette technique oblige l'émetteur à connaître tous les destinataires, elle surcharge le réseau. L'adressage de groupe (*multicast*) autorise l'émission d'un seul message qui ne sera dupliqué que par les nœuds ayant des clients abonnés à cette adresse dite de *multicast* (figure 11.10).

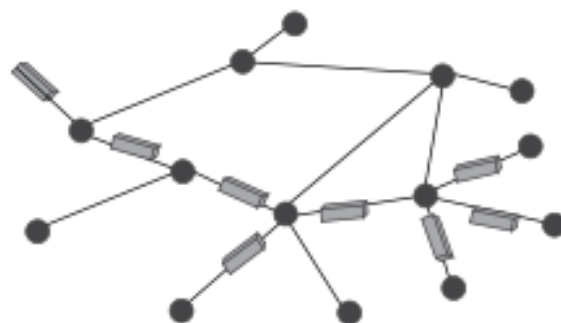


Figure 11.10 Principe du routage multicast.

□ Routage par inondation (de 1 vers tous)

Dans le routage par inondation, chaque nœud envoie le message sur toutes ses lignes de sortie, sauf celle d'arrivée du message (figure 11.11). Pour éviter une surcharge du réseau, chaque message comporte un compteur de sauts. Le compteur est initialisé à l'émission (nombre de sauts autorisés) et décrémente par chaque nœud. Le message est détruit quand le compteur de sauts est à zéro. Pour éviter les bouclages, les messages sont numérotés, chaque nœud mémorise cet identifiant et détruit les messages déjà vus.

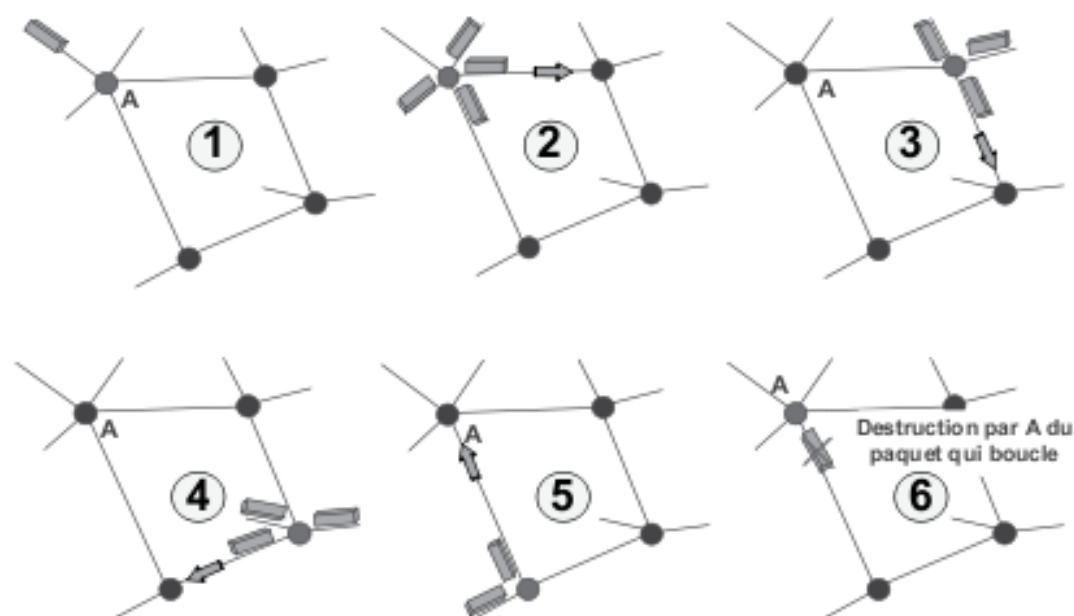
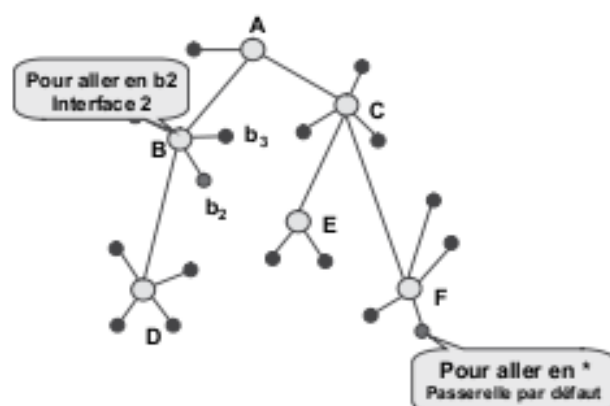


Figure 11.11 Le routage par inondation et la destruction des paquets qui bouclent.

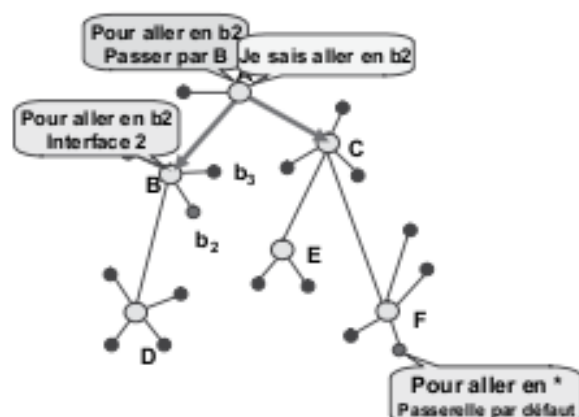
□ Routage par le chemin le plus court ou au moindre coût

Dans ce mode de routage, chaque nœud tient à jour des tables indiquant quel est le plus court (le meilleur) chemin pour atteindre le nœud destination. Chaque lien a un coût affecté ou calculé. À partir de ces informations de coût, chaque routeur détermine le chemin optimal pour joindre une destination. Ce coût ou métrique peut être exprimé en :

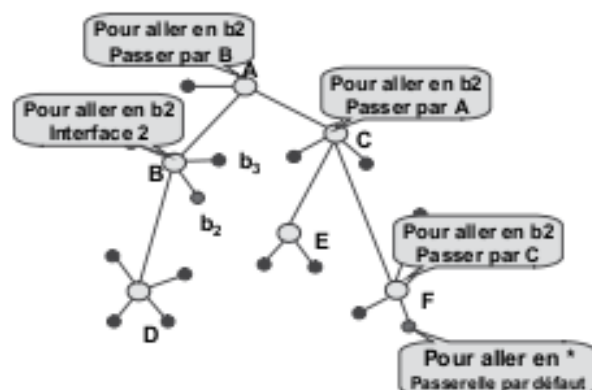
- ▶ nombre de sauts ;
- ▶ en temps de latence dans les files d'attente ;
- ▶ en délai de transmission ;
- ▶ fiabilité...



Les routeurs d'extrémité sont configurés manuellement par l'administrateur.



Chaque routeur diffuse à ses voisins sa base de connaissance, il sait comment joindre X.



Après un certain temps (temps de convergence) l'ensemble des routeurs du réseau sait comment joindre une destination d'extrémité.

Figure 11.12 Principe des algorithmes à vecteur distance.

□ Le routage au moindre coût

Principe des algorithmes vecteur distance

Dans le routage vecteur distance ou routage de Bellman-Ford (*distance vector routing*), chaque nœud du réseau maintient une table de routage qui comporte une entrée pour chaque nœud du réseau et le coût pour joindre ce nœud. Périodiquement chaque nœud diffuse sa table de routage à ses voisins. Le nœud destinataire apprend ainsi les destinations que son voisin sait joindre.

De proche en proche, chaque nœud apprend la configuration du réseau et le coût des différents chemins. La convergence des différentes tables peut être assez longue. Les schémas de la figure 11.12 illustrent ce propos.

Principe des algorithmes dits à état des liens

Le principal défaut du routage vecteur distance provient du fait que les routeurs n'ont la connaissance d'un changement d'état du réseau que lorsque leur voisin le leur communique, ce qui peut être long. Pour pallier ce défaut, le routage à état des liens (*link state routing*) procède différemment :

La figure 11.13 illustre le principe d'établissement des tables des routages, après découverte des voisins, chaque nœud diffuse le coût des liens qui le rattachent à ses voisins. À partir de ces informations, chaque nœud construit une matrice dite matrice des coûts et détermine, à partir de celle-ci, la route de moindre coût pour joindre chaque point du réseau.

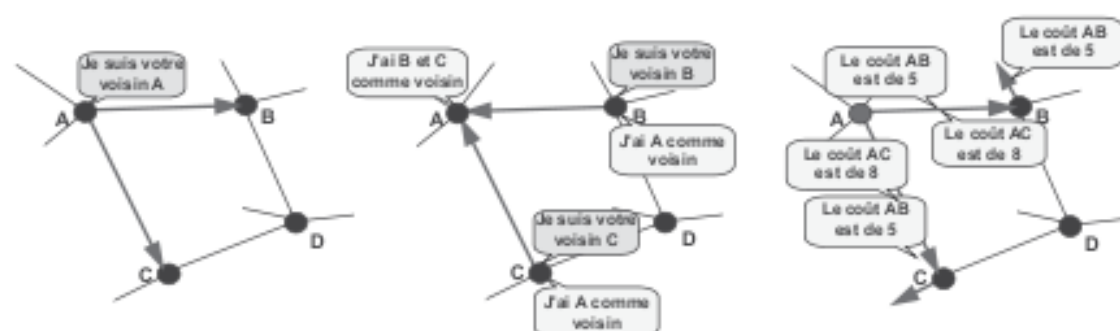


Figure 11.13 Principe des algorithmes à état des liens.

12 Le réseau IP

12.1 L'adressage dans IP

12.1.1 Le mode de mise en relation

Désirant alléger au maximum la couche inter-réseau, les concepteurs de TCP/IP n'ont spécifié qu'une couche réseau en mode non connecté (mode datagramme). Ce mode de mise en relation optimise l'utilisation des ressources réseaux mais n'assure ni contrôle d'erreur, ni contrôle de flux et de congestion. Au niveau du réseau, ces tâches peuvent éventuellement être assurées par le ou les sous-réseaux réels de transport (figure 12.1).

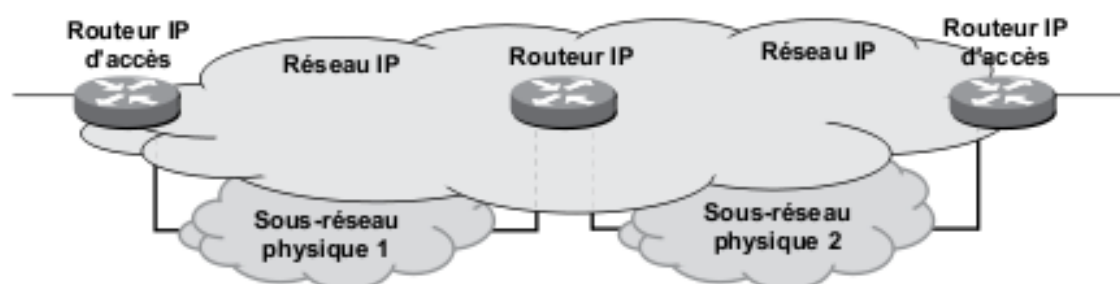


Figure 12.1 Le réseau IP (Internet Protocol).

Cependant, compte tenu qu'IP ignore la qualité de service offerte par ces sous-réseaux, la couche TCP devra pallier les insuffisances de la couche inter-réseau (*Internet Protocol*) en assurant le contrôle d'erreur, le contrôle de flux et de congestion. Cette approche, illustrée par la figure 12.2, reporte sur les systèmes d'extrémité des tâches normalement dévolues aux couches inférieures et en particulier le contrôle de congestion.

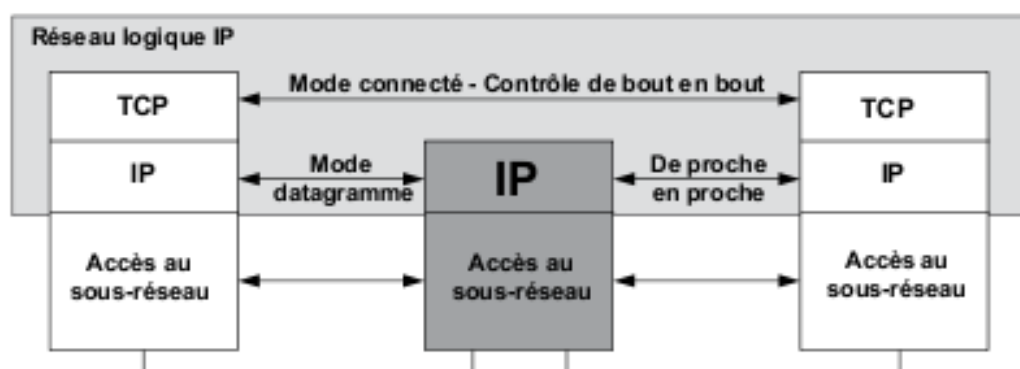


Figure 12.2 Le réseau logique IP et les modes de mise en relation.

12.2 L'adressage dans le réseau logique

Chaque machine (*Host*), raccordée au réseau logique IP, est, indépendamment de l'adressage physique utilisé dans le sous-réseau réel (figure 12.3), identifiée par un identifiant logique dénommé adresse IP. Le réseau logique IP masque le réseau physique. L'acheminement des données dans le réseau physique est réalisé à partir de l'adresse physique, les applications n'ont connaissance que de l'adresse logique IP, il est donc nécessaire d'établir une relation entre l'adresse logique et l'adresse physique correspondante (résolution d'adresses).

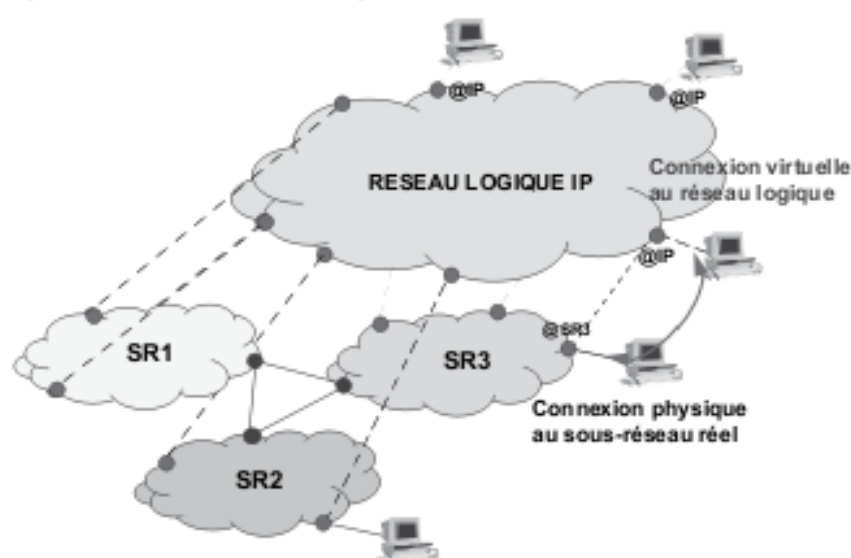


Figure 12.3 Nécessité d'une résolution d'adresses.

Enfin, le protocole IP doit assurer le routage dans le réseau logique IP. À cet effet, il doit pouvoir identifier le réseau logique IP (**Net_ID**) et, dans ce réseau, la machine cible (**Host_ID**). L'adressage logique IP ne comporte que ces deux informations (figure 12.4).

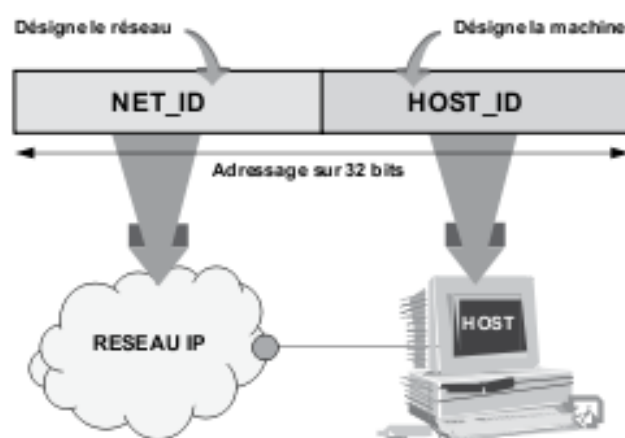


Figure 12.4 L'adressage dans le réseau logique IP.

12.3 Les techniques d'adressage dans un réseau IP

12.3.1 Les classes d'adressage

À chaque interface d'un système connecté à un réseau IP est assignée une adresse IP ou numéro IP. Cependant, l'adressage IP est un adressage à plat, il n'est pas possible, à partir de l'adresse IP, de déterminer la localisation géographique du réseau logique IP.

Limitée à 4 octets (32 bits), on représente l'adresse IPv4 par quatre valeurs décimales séparées par un point, la notation est dite décimale pointée (*Dotted-decimal notation*). Afin d'assurer une meilleure utilisation de l'espace d'adressage et d'adapter celui-ci à la taille et au besoin de chaque organisation, il a été introduit une modularité dans la répartition des octets entre l'identifiant réseau (**Net_ID**) et l'identifiant machine (**Host_ID**). Ainsi, cinq classes d'adresse (figure 12.5) ont été définies. Les premiers bits du champ adresse réseau (ID réseau ou **Net_ID**) permettent de distinguer la classe d'adressage.

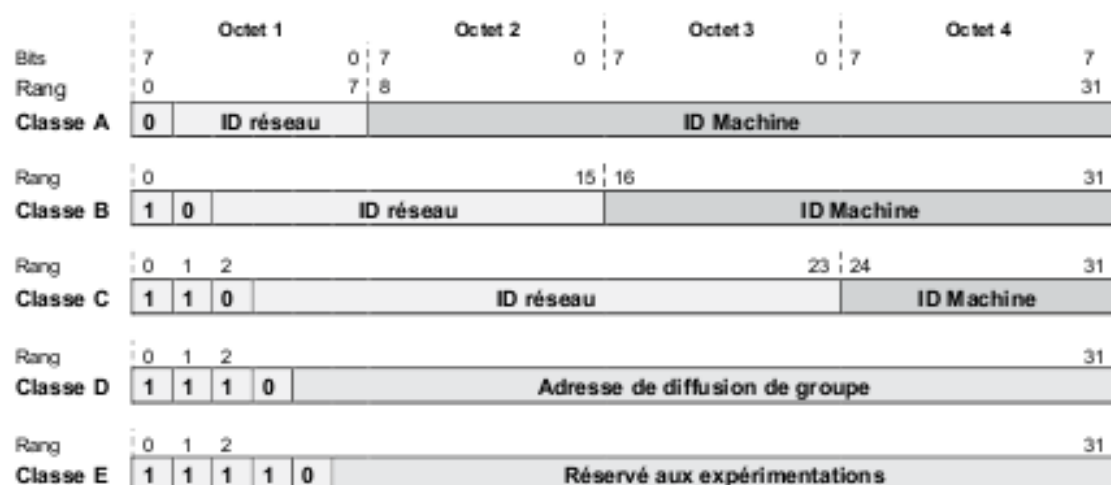


Figure 12.5 Les classes d'adresse IP.

Les adresses de classe A s'étendent de 1.0.0.1 à 126.255.255.254. Elles permettent d'adresser 126 réseaux¹ ($2^7 - 2$) et plus de 16 millions de machines² ($2^{24} - 2$, soit 16 777 214).

Les adresses de classe B vont de 128.0.0.1 à 191.255.255.254, ce qui correspond à 16 384 réseaux de 65 534 machines. Cette classe est la plus utilisée et les adresses sont aujourd'hui épuisées.

La classe C couvre les adresses 192.0.0.1 à 223.255.255.254, elle adresse 2 097 152 réseaux de 254 machines.

Les adresses de la classe D sont utilisées pour la diffusion (*Multicast*) vers les machines d'un même groupe. Elles vont de 224.0.0.0 à 239.255.255.255. Ce groupe peut être un ensemble de machines, mais aussi un ensemble de routeurs (diffusion des tables de routage). Tous les systèmes ne supportent pas les adresses de *multicast*.

Enfin, les adresses de la classe E sont réservées aux expérimentations.

12.3.2 Les adresses spéciales

Toute machine d'un réseau IP est identifiée par le couple <Net_ID><Host_ID>. Certaines valeurs de ces champs ont une signification particulière.

1 Les valeurs 0 et 127 sont réservées.

2 Les valeurs du champ Host_ID 0 et 255 ont une signification spécifique.

C'est ainsi que l'adresse $\langle \text{Net_ID} \rangle \langle 0 \rangle^1$, où tous les bits du champ Host_ID sont à zéro, désigne le réseau lui-même².

Certaines machines n'ont pas la possibilité de mémoriser une adresse IP. Lors du lancement de cette machine, le système émet une requête pour se voir attribuer une adresse IP (Protocole **RARP**, *Reverse Address Resolution Protocol*). Durant cette phase d'initialisation, la machine utilise l'adresse 0.0.0.0. Cette adresse ne peut donc pas être affectée à une machine particulière.

La machine elle-même ou machine locale peut être auto-adressée avec une adresse de la forme 127. x. x. x ; cette adresse dite de boucle locale (*loopback* ou encore *localhost*) est utilisée lors de tests de la machine ou de programmes applicatifs.

Lorsqu'une machine veut diffuser un message, elle peut, si le message ne s'adresse qu'à un ensemble de machines particulières, utiliser une adresse de *multicast* dite aussi de diffusion restreinte ou réduite. Si le message doit être adressé à toutes les machines, elle utilisera alors une adresse dite de diffusion générale. On distingue deux types d'adresses de diffusion générale :

- ▶ L'adresse 255.255.255.255 est utilisée pour envoyer un message à toutes les machines du même segment de réseau. La diffusion est limitée aux seules machines de ce segment, le datagramme n'est pas relayé sur d'autres réseaux. L'adresse 255.255.255.255 est dite adresse de **diffusion générale** ou **limitée** ;
- ▶ Si une machine veut s'adresser à toutes les machines d'un autre réseau, elle utilisera une adresse du type $\langle \text{Net_ID} \rangle \langle 1 \rangle$, tous les bits à 1 du champ Host_ID identifient toutes les machines du réseau $\langle \text{Net_ID} \rangle \langle 0 \rangle$. Ce message de diffusion est relayé de réseau en réseau pour

1 Pour simplifier et généraliser l'écriture des adresses, nous adopterons la convention d'écriture suivante : $\langle 0 \rangle$, tous les bits du champ concernés sont à zéro ; $\langle 1 \rangle$, tous les bits du champ concerné sont à 1.

2 En fait, cette adresse correspond à l'adresse de broadcast du système UNIX BSD version 4.2. Son emploi, comme adresse unicast, est possible mais est à déconseiller. Il est préférable d'en réserver l'utilisation à la désignation du réseau. Dans cet ouvrage, nous nous conformerons à cet usage.

atteindre le réseau destinataire. L'adresse est dite de **diffusion dirigée**. Ce principe est illustré figure 12.6.

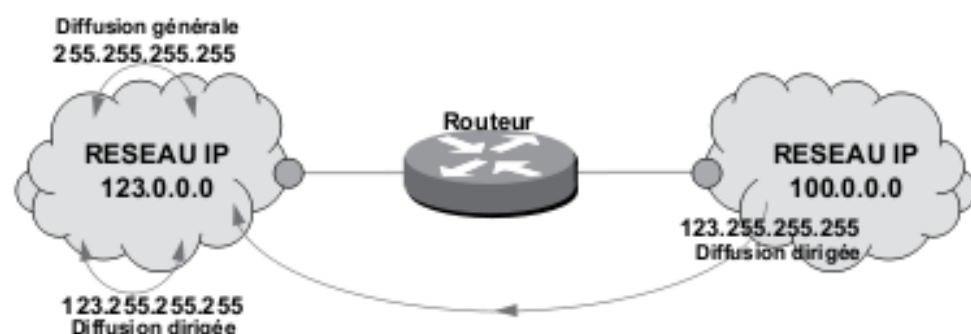


Figure 12.6 Les différentes adresses de diffusion.

Ainsi :

- ▶ l'adresse 123.0.0.0 désigne le réseau d'identifiant 123 ;
- ▶ l'adresse 123.0.0.18 désigne la machine 18 du réseau 123 ;
- ▶ l'adresse 123.255.255.255 est l'adresse de diffusion dirigée utilisée pour envoyer un message à toutes les machines du réseau 123 ; ce type d'adressage notamment utilisé, dans les réseaux locaux par certains protocoles de gestion (IEEE 802.1...) ;
- ▶ l'adresse 255.255.255.255 limite la diffusion aux seules machines du même segment de réseau que la machine source.

12.3.3 Les adresses publiques et les adresses privées

Pour permettre l'interconnexion des réseaux, il faut garantir l'unicité des adresses. C'est l'une des attributions de l'IANA (*Internet Assigned Numbers Authority*) qui attribue¹ à chaque réseau un identifiant unique. Certaines entreprises (organisations) disposent de leur propre réseau (réseau privé) et n'ont aucun besoin d'interconnexion vers l'extérieur, il est alors possible d'utiliser n'importe quelle adresse IP. Les adresses utilisées dans ce cas sont dites « illégales ». Par opposition, une adresse attribuée par l'IANA est dite publique ou légale.

¹ Depuis février 2011, l'IANA n'a plus d'adresses IPv4 disponibles.

Afin de prévenir, dans les réseaux privés, une éventuelle utilisation anarchique des adresses, il a été envisagé de réserver des plages d'adresses à ces réseaux (tableau 12.1). Ces adresses ne sont pas routables sur le réseau Internet. Elles sont réservées à un usage privé (RFC 1918), d'où l'appellation d'adresses privées.

Tableau 12.1 Les adresses privées (RFC 1918).

Classe	Début de la plage	Fin de la plage	Nombre de réseaux
A	10.0.0.0		1
B	172.16.0.0	172.31.0.0	16
C	192.168.0.0	192.168.255.0	256

Que faire si un réseau utilisant des adresses de type privé a soudainement des besoins d'accès à un réseau public ? Deux solutions sont envisageables :

- renuméroter toutes les stations avec des adresses publiques, non envisageable dans un grand réseau ;
- réaliser une conversion d'adresses (**NAT**, *Network Address Translator*), c'est-à-dire mettre en œuvre un mécanisme qui établit une correspondance entre une adresse privée et une adresse publique.

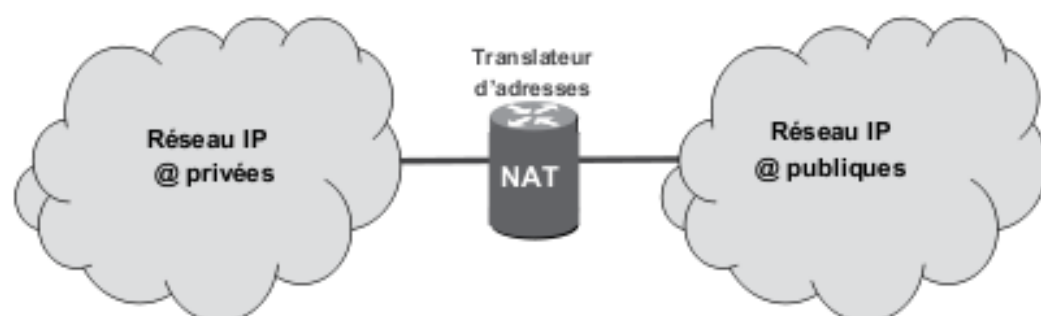


Figure 12.7 Le NAT est l'interface entre un réseau privé et un réseau public.

Pour des raisons évidentes de commodité, seule la seconde solution est généralement adoptée. C'est la passerelle d'accès au réseau public (routeur) qui réalise la translation d'adresses (figure 12.7). La traduction peut être statique, dans ce cas la table de correspondance est renseignée par l'administrateur du réseau, à une adresse privée on fait correspondre

une adresse publique. Cette approche limite les possibilités de connexion vers l'extérieur. La traduction peut aussi être dynamique, la mise en correspondance adresse privée/adresse publique est établie au moment du besoin d'interconnexion de la machine du réseau privé. Les adresses publiques peuvent alors être partagées par l'ensemble des machines du réseau privé. La traduction dynamique permet de n'utiliser qu'un nombre restreint d'adresses publiques voire une seule pour N machines.

12.3.4 Notion de sous-réseau : le *subnetting* (RFC 950)

■ Nécessité de définir des sous-réseaux

Une entreprise peut disposer d'un réseau découpé en N sous-réseaux (locaux ou territorialement dispersés). Pour assurer l'acheminement dans le réseau global de l'entreprise, il est nécessaire de pouvoir identifier ces sous-réseaux. L'adresse réseau IP ne contenant qu'un seul identifiant se révèle insuffisante pour assurer cette tâche. Aussi, en 1984, a-t-on introduit la notion d'identifiant de sous-réseaux (*subnetting*).

Dans cette technique, tout nœud du réseau logique est différencié et localisé par un unique identifiant du réseau logique <Net_ID> (le réseau global de l'entreprise), un identifiant du site local ou de sous-réseau appelé <SubNet_ID> et enfin un numéro de nœud ou *host* (figure 12.8). La structure originelle de l'adressage IP ne prend pas en compte cette structure d'adresse. Aussi, compte tenu que généralement tous les bits du champ <Host_ID> ne sont pas utilisés pour numérotter les machines, il suffit donc d'en prélever quelques-uns pour constituer le champ <SubNet_ID> d'identification du sous-réseau. La taille de ce champ sera déterminée en fonction du nombre de sous-réseaux à distinguer.

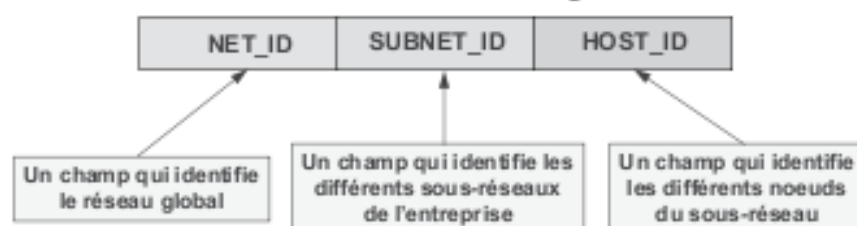


Figure 12.8 La technique du subnetting décompose l'adresse IP en trois champs.

En principe, l'acheminement est réalisé à partir du champ <Net_ID> dont la taille, dépendant de la classe d'adressage, est connue de chaque routeur. L'utilisation d'un identifiant réseau de longueur variable nécessite de fournir à chaque nœud du réseau les informations qui lui permettent de distinguer les bits d'adressage à prendre en compte pour définir l'acheminement dans le réseau de ceux qui n'ont qu'une signification locale. Ces informations sont fournies, à chaque nœud adressé, sous forme d'un champ de bits à 1 appelé **masque de sous-réseau** (figure 12.9).

Il existe deux méthodes d'écriture des masques de sous-réseaux, qui sont équivalentes :

- ▶ réseau : 10.0.0.0, masque de sous-réseau 255.255.240.0 ;
- ▶ ou plus simplement 10.0.0.0/20, le préfixe 20 indique la longueur en bits du masque de sous-réseau (longueur du préfixe réseau ou simplement préfixe).

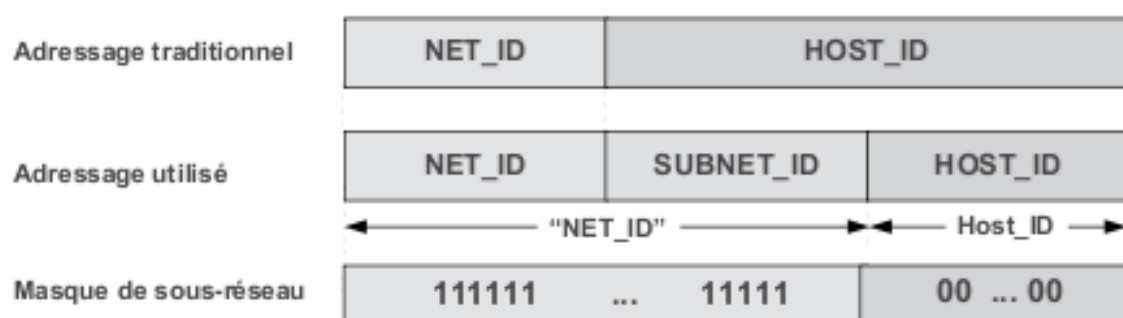


Figure 12.9 Principe du masque de sous-réseau.

■ Utilisation du masque de sous-réseau

Lorsqu'une station émet un datagramme à destination d'une autre station, pour déterminer si la machine cible est localisée sur le même sous-réseau, la machine source réalise un « ET » logique entre les bits de sa propre adresse (adresse source) et ceux du masque de sous-réseau, elle procède de même avec l'adresse destination. Si le résultat donne une valeur identique, les deux machines sont sur le même sous-réseau, le datagramme y est diffusé ; sinon le datagramme est adressé à la passerelle par défaut (figure 12.10), charge à elle de trouver le sous-réseau destinataire.

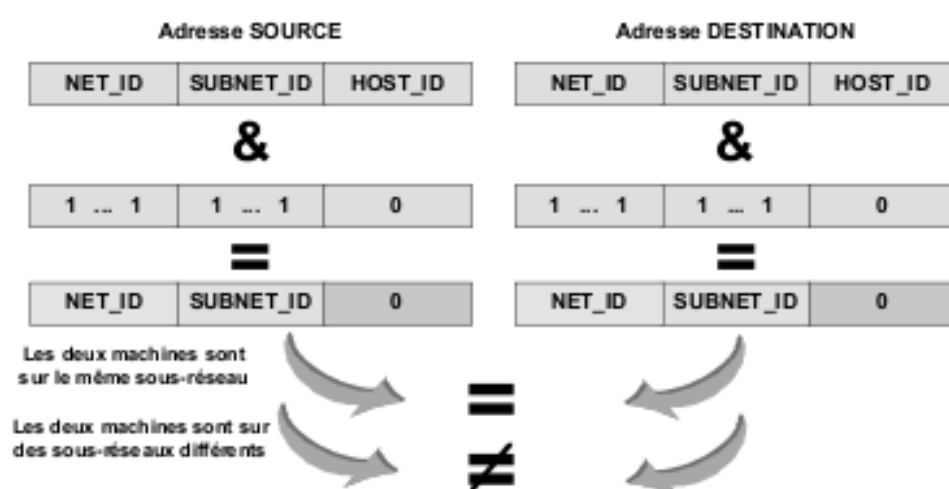


Figure 12.10 Détermination du sous-réseau cible à l'aide du masque de sous-réseau.

12.3.5 L'adressage géographique ou CIDR

Une adresse IP désigne une organisation, elle ne permet pas d'en déterminer la localisation, c'est un adressage à plat. Dans ces conditions, chaque routeur du réseau Internet doit tenir à jour la liste de toutes les adresses attribuées (Net_ID) et la route à suivre. Cet encombrement des tables de routage pénalise les performances ; il a conduit, lors de la recherche de solutions pour pallier la prévisible pénurie d'adresses IPv4 à mettre en œuvre un mécanisme d'affectation géographique des adresses de classe C non attribuées (RFC 2050), quelques exemples sont donnés dans le tableau 12.2.

Tableau 12.2 L'allocation géographique des adresses de classe C.

Plage d'adresses	Zone d'affectation
192-193	Divers (adresses déjà attribuées)
194-195	Europe (65 536 réseaux)
196-197	Divers
198-199	Amérique du Nord
200-201	Amérique centrale et du Sud
202-203	Pacifique
204-205	Divers
206-207	Divers

D'autre part, il a été décidé de n'attribuer qu'exceptionnellement les adresses de classes B restantes et d'attribuer en lieu et place des adresses contiguës de classe C en leur faisant correspondre une seule entrée dans les tables de routage.

Ainsi, pour l'Europe, les adresses 194 et 195 ont les sept premiers bits identiques. Il suffit donc d'indiquer aux routeurs que le champ <Net_ID> à prendre en compte est de 7 bits et non de considérer ces adresses comme des adresses de classes C. Une seule entrée suffit alors dans la table de routage. Le nombre de bits servant à coder la partie commune, ou préfixe d'adresse, est représenté à la fin de l'écriture de l'adresse comme suit : 194.0.0.0/7, ainsi cette adresse désigne tous les sous-réseaux européens. Cette technique, issue de celle du masque de sous-réseau, porte le nom de *supernetting* ou routage interdomaine sans tenir compte de la classe d'adressage (**CIDR**, *Classless InterDomain Routing*).

12.4 La structure du datagramme IP

Un datagramme IP (figure 13.3) peut contenir (champ données) un segment TCP, un datagramme UDP, un message ICMP, ARP, RARP ou encore OSPF. Les différents champs de l'en-tête IP sont alignés sur des mots de 32 bits (figure 12.11), si aucune option n'est invoquée cet en-tête comporte 20 octets (cinq mots de 32 bits). Les bits sont représentés dans l'ordre d'émission sur le support (bits de poids fort devant, dit aussi *big endian*).

Le numéro de version sur 4 bits (**VER**) permet d'identifier la version d'IP en cours et par conséquent le format du datagramme. La présence de cette information autorise la cohabitation de plusieurs versions de protocole dans les systèmes intermédiaires (IPv4, IPv6...).

Le champ longueur d'en-tête sur 4 bits (**IHL**, *Internet Head Length*) indique, en multiple de mots de 32 bits, la longueur de l'en-tête. Lorsqu'aucune option n'est invoquée, ce champ vaut cinq (20 octets).

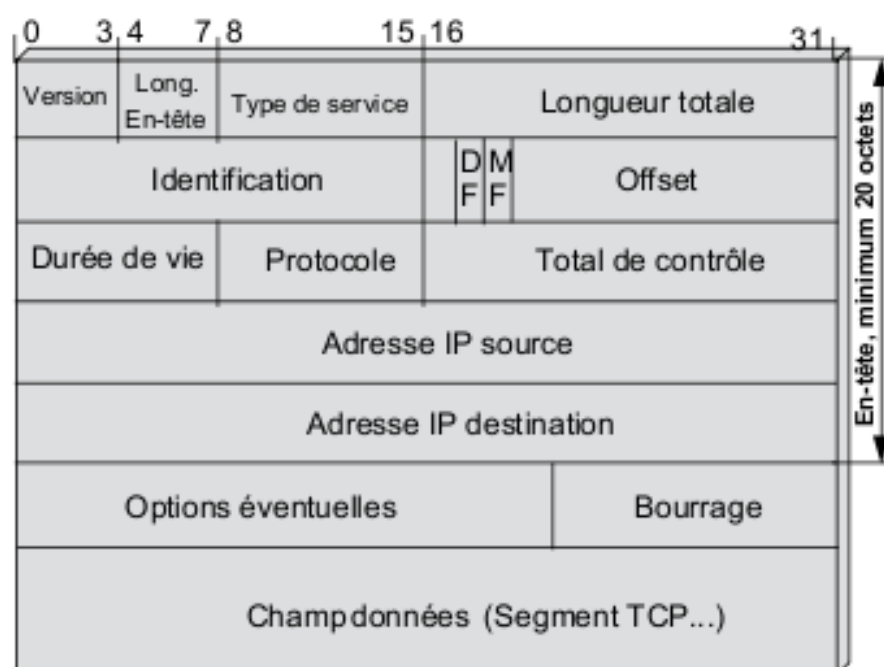


Figure 12.11 La structure du datagramme IP.

Le champ type de service sur 8 bits (**TOS**, *Type Of Service*) spécifie, à la passerelle inter-réseau, le type d'acheminement attendu. Le RFC 791 a défini huit niveaux de priorité et quatre critères d'acheminement (figure 12.12), mais la plupart des systèmes intermédiaires n'ont pas la possibilité de traiter ces informations. Définies alors que les flux dans le réseau étaient tous de même nature (texte), les priorités du champ TOS sont inadaptées au transport des flux multimédias. Le RFC 1812 (*IP Precedence*) a modifié l'interprétation du champ TOS, en redéfinissant les trois premiers bits (figure 13.1). Cette nouvelle interprétation a été jugée insatisfaisante pour garantir aux différents flux IP une certaine qualité de service et un acheminement qui la respecte. Il est apparu nécessaire de différencier plus finement les flux pour, en cas de congestion dans le réseau, mettre en œuvre une politique d'acheminement et d'élimination des datagrammes en relation directe avec la qualité de service exigée par les différents flux IP. Le RFC 2474 (*Differentiated Services* ou *DiffServ*) introduit la notion de classes de service et définit, pour chaque classe, en cas de congestion, différentes politiques d'élimination des datagrammes (figure 12.12).

Champ TOS	Priorité	D Délai	T Débit	R Fiabilité	C Coût	Res
IP Precedence	Precedence	D Délai	T Débit	R Fiabilité	CU	
DiffServ	DSCP (Differentiated Services Code Point)					CU

Figure 12.12 Comparaison des différentes interprétations du champ TOS.

Le champ DS (*DiffServ*) de la figure 12.12 comporte deux sous-champs, DSCP et CU. Le sous-champ **DSCP** (*Differentiated Service Code Point*), sur 6 bits, autorise 64 classes de trafic réparties en trois grandes familles :

- ▶ *Expedited Forwarding* (EF, RFC 2598), défini spécifiquement pour les applications temps réel, minimise le temps de latence dans le réseau. Ce service est dénommé *Premium Service* ;
- ▶ *Assured Forwarding* (AF, RFC 2597) ou *Olympic Service* comprend quatre classes, elles-mêmes subdivisées. À chaque classe correspond une priorité différente avec une garantie de bande passante. Des mécanismes spécifiques permettent l'élimination de datagrammes en cas de congestion. Ces différentes classes sont généralement désignées sous les termes de : services Platinum, Gold, Silver ou Bronze ;
- ▶ *Best effort*, aucun traitement spécifique n'est réalisé. Les datagrammes sont transmis pour le mieux.

Le sous-champ CU (*Currently Unused*) sur 2 bits est non utilisé dans la version standard de TCP/IP. Dans la version de TCP/ECN, ce champ dit ECN (*Explicit Congestion Notification*) indique à l'entité distante que l'émetteur implémente un mécanisme spécifique de contrôle de congestion dit ECN.

Le champ **Longueur totale** (**LEN**, *total LENGTH field*) sur 16 bits indique la longueur totale, en octets, du datagramme en-tête compris. La longueur maximale est de 65 536 octets. Cette longueur maximale est toute théorique, elle dépend évidemment des capacités du réseau qui pour un réseau IP ne saurait être en dessous de 576 octets (MTU minimale d'un réseau IPv4).

Le champ **Identification (ID)** sur 16 bits : la valeur du champ ID, attribuée par la source, est générée de manière aléatoire par un algorithme initialisé par l'heure système. En cas de fragmentation, l'ID est recopiée par les systèmes intermédiaires dans tous les fragments du datagramme d'origine. L'ID permet à l'hôte destinataire d'identifier (N° identification et adresse IP) les différents fragments d'un même datagramme, il facilite ainsi le réassemblage.

Le champ suivant est composé de 3 bits dont le premier n'est pas utilisé. Le bit suivant dit bit **DF** (*Don't Fragment*, ne pas fragmenter), s'il est positionné à 1, demande aux systèmes intermédiaires de ne pas fragmenter le datagramme. Ce bit est utilisé, par exemple, quand le système d'extrémité est incapable de réassembler les différents fragments. Le système intermédiaire qui reçoit un tel datagramme doit soit le router dans sa totalité sur un sous-réseau où la MTU est compatible, soit le détruire. En cas de destruction, il en avertit la source par un message ICMP (*Internet Control Message Protocol*). Enfin, le bit **MF** (*More Fragment*) est positionné à 1 dans tous les fragments d'un même datagramme d'origine pour indiquer qu'un fragment suit. Il est à 0 dans le dernier fragment ou lorsqu'un datagramme n'a pas subi de fragmentation.

En cas de fragmentation, le champ *Offset* (13 bits) indique, en multiple de 8 octets, la position du premier bit du fragment dans le datagramme d'origine. En conséquence, tous les fragments, sauf le dernier, doivent avoir une longueur multiple de huit (voir § 4.5).

Le champ **Durée de vie (TTL, Time To Live)** sur 8 bits exprime en seconde, la durée de vie d'un datagramme. Cette valeur est décrétementée à chaque passage à travers une passerelle. Un datagramme avec un TTL à zéro est détruit. Aucune estampille de temps ne figurant dans l'en-tête IP, les passerelles (routeur) n'ont pas la possibilité de mesurer le temps écoulé, elles se contentent alors de décrétement ce champ d'une unité, ce qui revient à définir non pas une durée de vie, mais un nombre de sauts. Le TTL est généralement initialisé à 32 voire 64 sauts.

Le champ **protocole (Protocol)**, 8 bits, indique à IP l'origine du champ données (protocole transporté, RFC 1700, *Assigned numbers*). Ce champ permet le multiplexage de flux.

Le champ **total de contrôle** (*Header Checksum*), 16 bits, n'est calculé que sur l'en-tête IP. Le total de contrôle est le complément à 1 de la somme en complément à 1 des données de l'en-tête découpées en mots de 16 bits. Le total de contrôle est recalculé par chaque système intermédiaire (modification du champ TTL et fragmentation éventuelle).

Le champ **option**, de longueur variable, est codé : code option (type), longueur, valeur. Coûteux en termes de traitement par les passerelles, le champ option est peu, voire pas utilisé.

12.5 Le contrôle de la fragmentation sous IP

La fragmentation d'un datagramme IP est contrôlée par les champs : longueur totale (LEN), offset (Offs) dans le segment, et le bit MF (*More Fragment*) du datagramme IP. Le champ offset indique, en multiples de 8 octets, la position du fragment dans le datagramme initial. Le fragment ainsi constitué ne peut avoir pour longueur que le multiple de huit le plus proche de la MTU, sauf pour le dernier fragment.

Ainsi, pour une MTU de 128 octets¹, la charge utile (niveau IP) ne peut être, au maximum, que de 108 octets ($128 - 20$ d'en-tête IP), soit une taille effective de 104 octets (13×8), si l'on tient compte de l'en-tête TCP (20 octets), la charge du premier fragment n'est que de 84 octets. La figure 12.13 illustre la fragmentation d'un datagramme IP de 576 octets. Pour faciliter la compréhension, dans la figure 12.13, les valeurs des champs *Len* et *Offset* sont exprimées en décimal.

¹ Valeur courante de la MTU, lors d'un transit dans un réseau X.25.

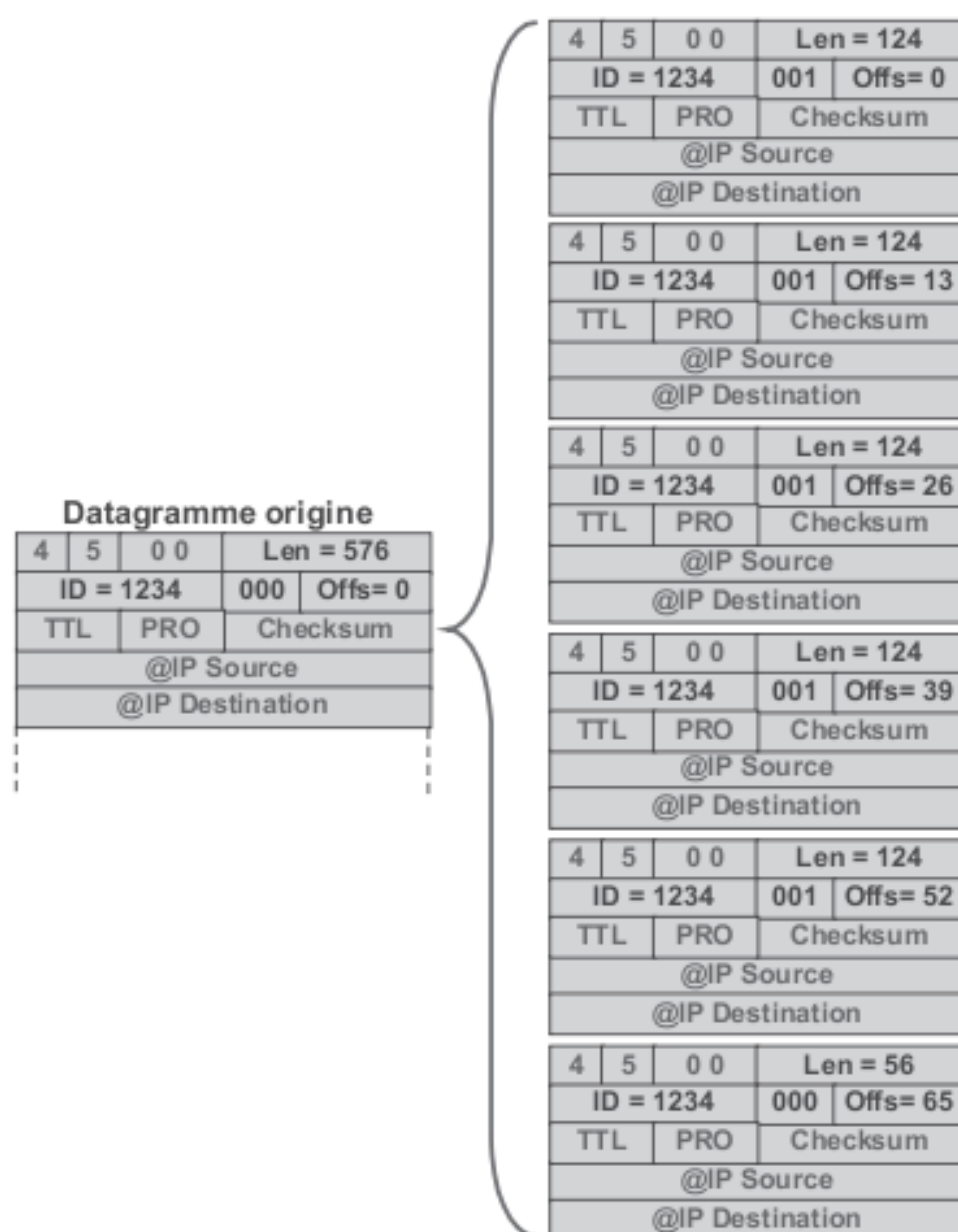


Figure 12.13 La fragmentation d'un datagramme IP.

Pour assurer le réassemblage, IP doit attendre l'arrivée de tous les fragments. Les opérations de fragmentation et de réassemblage sont coûteuses en termes de puissance de calcul et de mémoire. De plus, en mode datagramme, la perte d'un seul fragment provoque une reprise par la couche TCP de la totalité du segment fragmenté.

13

D'IPv4 à IPv6

13.1 Les lacunes d'IPv4

IPng (*next generation*) ou IPv6 répond au besoin d'évolution de la communauté Internet et comble les faiblesses d'IPv4. La plus connue concerne l'espace d'adressage dont la structure à plat (*Net_ID*) est responsable de l'explosion des tables d'adressage et dont la structure figée est responsable de la pénurie d'adresses. En faisant disparaître la notion de classe d'adresses, en autorisant l'agrégation d'adresses de réseaux contigus en un seul préfixe réseau et en organisant une affectation géographique des adresses, le CIDR a partiellement résolu ces problèmes. L'utilisation d'un adressage privé associé à la translation d'adresses (NAT) a repoussé les limites de l'espace d'adressage mais en pénalisant les performances.

Enfin, l'arrivée de nouvelles applications comme le multimédia et le besoin de services sécurisés ont motivé l'étude d'un nouveau protocole permettant d'augmenter l'espace d'adressage tout en conservant les grands principes qui ont fait le succès du protocole IP. Les principales caractéristiques d'IPv6 sont :

- ▶ adressage étendu (128 bits au lieu de 32) ;
- ▶ en-tête simplifié autorisant un routage plus efficace ;
- ▶ sécurité accrue en incluant des mécanismes d'authentification, de cryptographie et en garantissant l'intégrité des données ;
- ▶ implémentation d'un mécanisme de découverte de la MTU optimale. La fragmentation n'est plus réalisée dans le réseau mais par le nœud source ;
- ▶ suppression du champ *Checksum*, ce qui allège le travail des routeurs intermédiaires ;

- ▶ amélioration des aspects de diffusion (*Multicast*) ;
- ▶ intégration de fonctions d'autoconfiguration et de renumérotation.

13.2 L'adressage dans IPv6

13.2.1 Généralités

L'adressage dans IPv6 a été porté à 128 bits (2^{128} adresses soit plusieurs milliards d'adresses par mètre carré de surface terrestre : $6,65 \cdot 10^{23} @/m^2$).

Dans un système de réseaux interconnectés, seul un adressage hiérarchique permet l'allègement des tables de routage. Cependant, dans une communauté aussi vaste que celle d'Internet, l'adressage hiérarchique géographique devient vite sans signification, aussi entre l'adressage à plat non significatif d'IPv4 et l'adressage hiérarchique, tel que celui d'X.121, un compromis a été réalisé. L'adressage IPv6 comporte quatre champs (figure 13.1).

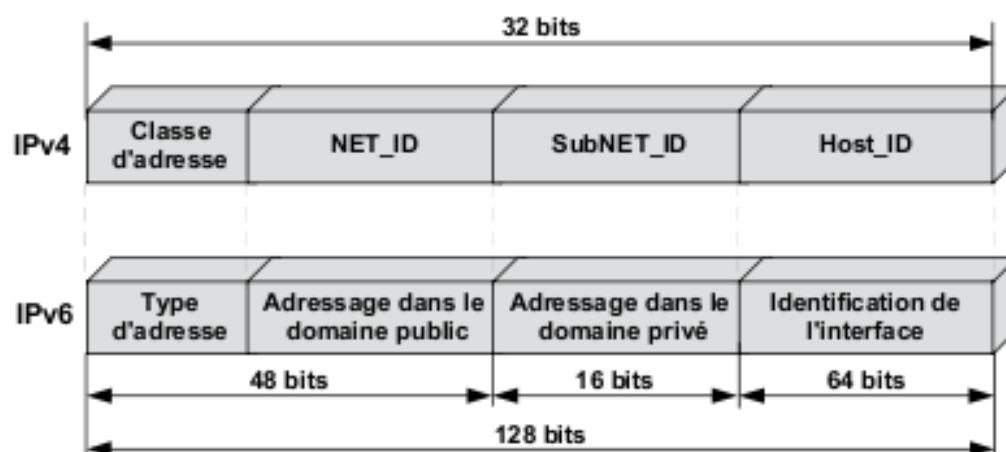


Figure 13.1 Principe de l'adressage IPv6.

Le premier, sur 48 bits, est une agrégation hiérarchique de préfixes décrivant la connectivité du site, ce champ est désigné sous le terme topologie publique, il est attribué par l'opérateur auquel on est raccordé. L'affectation d'adresses par l'opérateur change radicalement la nature de l'adressage, on n'est plus propriétaire de ses adresses ; si on change d'opérateur, on change d'adresse !

Le second, sur 16 bits, décrit la topologie locale du site, enfin le dernier, sur 64 bits, identifie de manière unique une interface. Cet adressage est dénommé adressage agrégé ou *Aggregatable Global Unicast Address Format*.

13.2.2 La notation IPv6

Une notation hexadécimale, sur 16 bits séparés par deux points « : », remplace la notation décimale pointée d'IPv4. L'adresse passe de 32 à 128 bits, huit mots de 16 bits. Ainsi, une adresse IPv6 s'écrit :

FE0C :DA98 :0 :0 :0 :0 :5645 :376E

La notation peut être simplifiée en remplaçant une succession de 0 par « :: », l'abréviation « :: » ne pouvant être utilisée qu'une seule fois. Ainsi, l'adresse précédente devient :

FE0C :DA98 ::5645:376E

IPv6 adopte une notation similaire à celle du CIDR, le champ préfixe étant désigné par un nombre représentant la longueur en bits du préfixe, l'écriture est donc de la forme : @IPv6/ longueur du préfixe en bits, soit par exemple :

FE0C:DA98/32

FE0C:DA98 :0:0/64

FE0C:DA98 ::/64

Une adresse IP peut être utilisée comme URL (*Uniform Resource Locators*), par exemple en IPv4 : `http://80.12.4.212:8080` où : 8080 désigne l'application distante (port). En IPv6, on obtiendrait une URL du type : `http:// FE0C:DA98 : 5645:3763:8080` ce qui introduit une ambiguïté, 8080 est le port destination où le dernier champ de l'adresse ? Aussi, le RFC 2732 propose d'écrire l'adresse IPv6 entre « [] », dans notre exemple cela donne :

`http ://[FE0C:DA98 ::5645:3763]:8080`

13.2.3 Les types d'adresse

Très pénalisante en termes de performance réseau, la notion de *broadcast* disparaît. Elle est remplacée par une généralisation des adresses *multicast*. IPv6 distingue trois types d'adresses :

- ▶ les adresses **unicast** (*one-to-one*) : une adresse *unicast* désigne une interface, elle peut être utilisée pour identifier un groupe d'interfaces lorsque ces interfaces constituent une agrégation de liens et doivent être vues comme une seule interface ;
- ▶ les adresses **multicast** (*one-to-any*) : ces adresses désignent un ensemble d'interfaces dont la localisation n'est pas nécessairement sur le même réseau physique. Un datagramme adressé à une adresse *multicast* est acheminé à toutes les interfaces du groupe ;
- ▶ les adresses **anycast** (*one-to-nearest*) : ces adresses introduites par IPv6 correspondent à une restriction des adresses de *multicast*. Elles désignent un ensemble d'interfaces partageant un même préfixe réseau. Cependant, lorsqu'un datagramme est adressé à une adresse *anycast*, il n'est délivré qu'à une seule interface du groupe, celle dont la métrique, au sens routage du terme, est la plus proche du nœud source.

13.2.4 L'adressage de transition d'IPv4 vers IPv6

La migration d'IPv4 vers IPv6 ne peut être réalisée que progressivement. Durant un temps assez important, les deux versions du protocole devront cohabiter (RFC 2893). Aussi, deux solutions ont été envisagées pour permettre d'utiliser les adresses IPv4 dans le domaine IPv6 (AF_INET6). La première (figure 13.2) dite « adresses IPv4 mappées » est une représentation interne des adresses afin de permettre à des programmes IPv6 de fonctionner sur un réseau IPv4, la communication se faisant d'une machine IPv4 à IPv4.

La seconde dite IPv4 compatible (usage déconseillé, obsolète RFC 4291), permet à deux machines IPv6 de communiquer à travers un réseau IPv4 (tunnel IPv4). Le datagramme IPv6 d'adresse « ::a.b.c.d » est encapsulé dans un datagramme IPv4 d'adresse « a.b.c.d ».

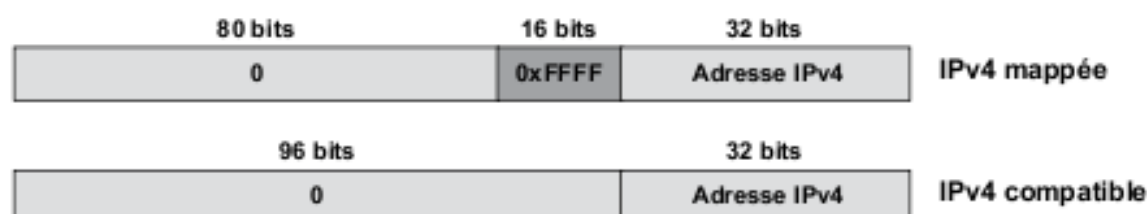


Figure 13.2 L'adressage IPv4/IPv6.

13.3 Le datagramme IPv6

13.3.1 La structure du datagramme

Pour améliorer le traitement des datagrammes dans les routeurs, la structure même du datagramme a été simplifiée en supprimant les champs devenus inutiles et en particulier le champ option. La suppression de ce dernier confère à l'en-tête du datagramme une longueur constante (40 octets). La figure 13.3 compare les datagrammes IPv4 et IPv6. Les champs IPv4 qui n'ont aucune correspondance dans le datagramme IPv6 ont leur label grisé et en italique.

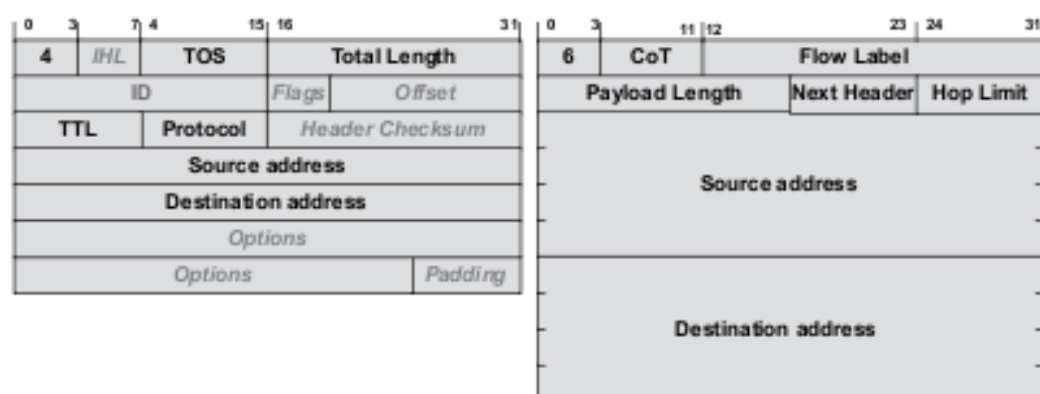


Figure 13.3 Les datagrammes IPv4 et IPv6.

Les données relatives à la fragmentation (ID, Flags, Offset) disparaissent de l'en-tête. IPv6 implémente un mécanisme de découverte de la MTU. Si elle est nécessaire, la fragmentation est réalisée par la source et le réassemblage par le destinataire. Ceci allège considérablement le travail des routeurs intermédiaires. Cependant, si la fragmentation par le réseau s'avère indispensable, une extension d'en-tête est prévue à cet effet. L'en-tête étant de longueur fixe, le champ IHL (longueur de l'en-tête) est devenu

inutile. De même, le calcul du total de contrôle a été supprimé, le mécanisme de contrôle de TCP incluant un pseudo en-tête IP est suffisant pour protéger les adresses.

Le champ TOS (*Type of Service*) trouve son équivalent en deux champs, un champ classe de trafic (**CoT**, *Class of Traffic*) sur 8 bits et une identification de flux (*Flow label*). Ce dernier champ contient un identifiant attribué initialement par la source, chaque routeur mémorise ce label et tente de router identiquement tous les datagrammes d'un même flux. Cette information va contribuer à la mise en place d'une véritable gestion de la qualité de service dans les réseaux IP et notamment faciliter le développement d'application de type voix. Le contexte ainsi créé est détruit sur temporisation (Timer d'inactivité). Le réseau ainsi défini est dénommé *soft-state*.

L'en-tête étant de longueur fixe, le champ longueur totale d'IPv4 (*Total length*) est remplacé par la taille des données transportées (*Payload length*). Enfin, le champ *Protocol* d'IPv4 est remplacé par l'indication du type de l'en-tête suivant (*Next header*). Si aucune option n'est invoquée, ce champ contient l'identification du protocole transporté, la figure 13.4 illustre le chaînage des en-têtes.

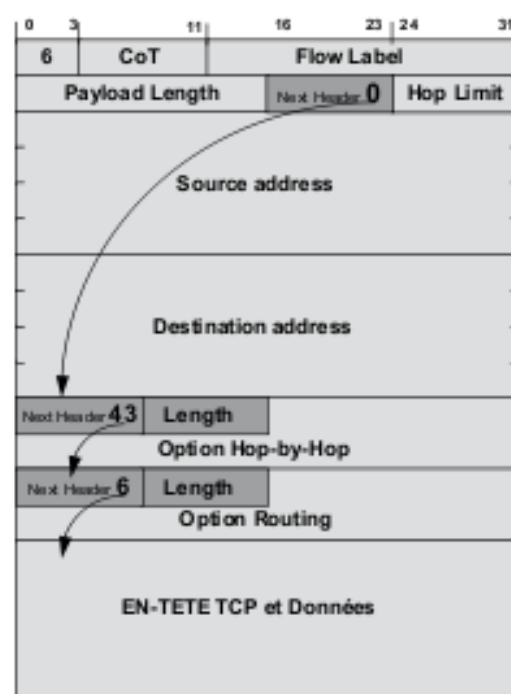


Figure 13.4 Principe du chaînage des en-têtes.

13.4 Conclusion

Le protocole IP, imaginé à une époque où seules des informations de type texte transitaient dans les réseaux, a su, en raison de sa simplicité originale, s'adapter et représente pratiquement aujourd'hui le seul protocole mis en œuvre dans les réseaux. La nouvelle version d'IP (*IP next generation* ou IPv6), apporte des améliorations sensibles surtout en matière de sécurité. Longtemps retardée, son arrivée dans les réseaux privés ne saurait tarder même si la cohabitation IPv4 et IPv6 devrait durer encore de nombreuses années.

5

Les protocoles de transport : TCP et UDP

14

Les mécanismes de base de TCP

S'appuyant sur un protocole réseau non fiable (*Best effort*), **TCP** (*Transmission Control Protocol*, RFC 793) est un protocole de niveau transport en mode connecté (figure 14.1). TCP garantit la délivrance des données en séquence, il en contrôle la validité et organise les éventuelles reprises sur erreur ; enfin, il effectue un contrôle de flux de bout en bout et met en œuvre un mécanisme de détection et de guérison de la congestion.

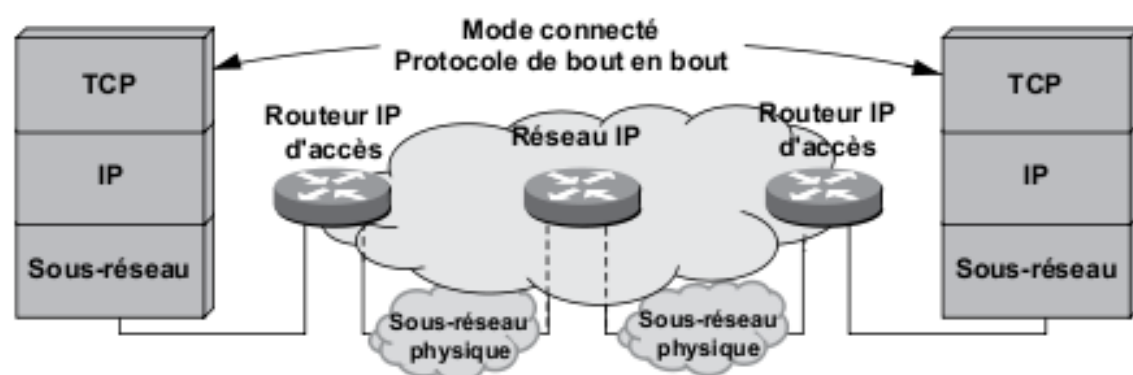


Figure 14.1 Transmission Control Protocol et IP.

14.1 La notion de connexion de transport

La connexion au niveau transport permet de voir, au niveau applicatif, le réseau comme un lien virtuel entre deux applications actives sur les systèmes d'extrémité. Ce « lien virtuel » correspond à l'identification des processus communicants distants. Cette connexion est complètement définie par l'association de (figure 14.2) :

{protocole, port destination, @IP destination, port source, @IP source}

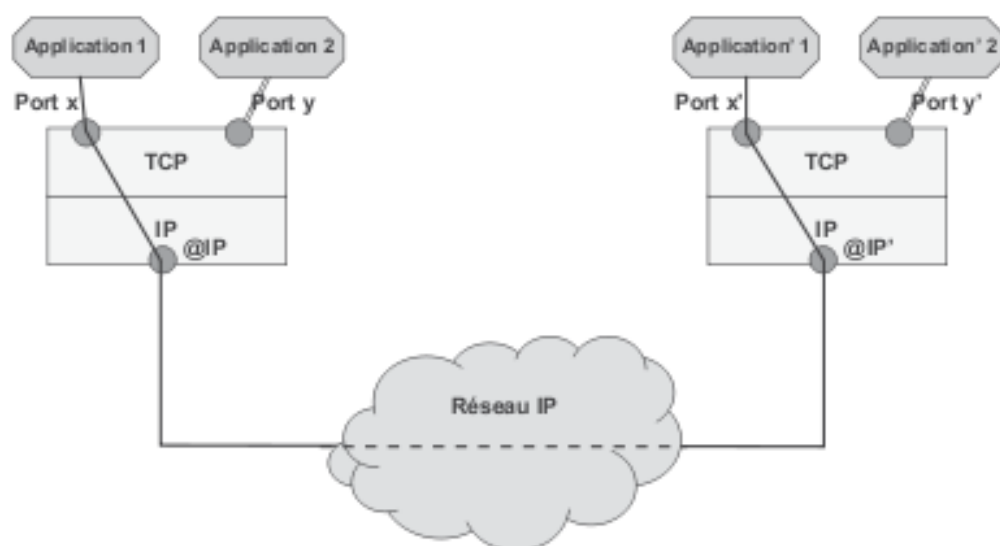


Figure 14.2 La connexion de transport.

Cette association est désignée sous le terme de *socket* (RFC 793). En fait, la notion de *socket* est beaucoup plus large ; dans le monde UNIX, un *socket* définit une interface de programmation ou **API**, *Application Programming Interface*. L'environnement Microsoft définit un concept identique : les WinSock.

Un même système peut exécuter plusieurs applications, il doit donc être capable d'établir et de distinguer plusieurs connexions de transport. C'est le concept de port qui, outre l'identification des applications, autorise le multiplexage des connexions de transport (figure 14.3).

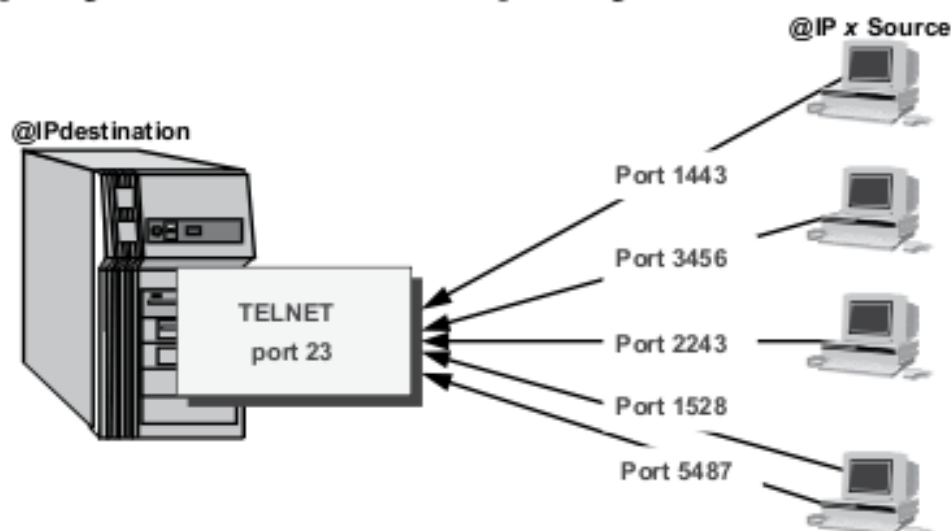


Figure 14.3 Le multiplexage des connexions de transport.

Le tableau 14.1 fournit quelques exemples de numéro de port tels qu'ils sont décrits dans le RFC 1700 (*Assigned numbers*). Les 1 024 premières valeurs sont réservées. Les ports identifiés par ces valeurs sont dits **référéncés** ou encore **ports bien connus** (*well known ports*), ces ports sont liés à des applications spécifiques.

Tableau 14.1 Exemples de numéros de ports référencés.

Nom du service	Numéro de port/ protocole	Alias	Commentaire
echo	7/tcp		
ftp-data	20/tcp		Protocole de transfert de fichiers
ftp	21/tcp		
telnet	23/tcp		Émulation de terminal
domain	53/udp	nameserver	Serveur de noms de domaine
x400	104/tcp		Courrier ISO
pop3	110/tcp	postoffice	Bureau de poste
nbssession	139/tcp	netbios_ssn	Session de service Netbios

Les ports référencés permettent ainsi, à une application dite **cliente**, de désigner une application sur le système distant. L'extrémité de la connexion cliente est identifiée par un numéro de port attribué dynamiquement par le processus appelant. Ce numéro de port est dit **port dynamique** ou **éphémère** (figure 14.4).

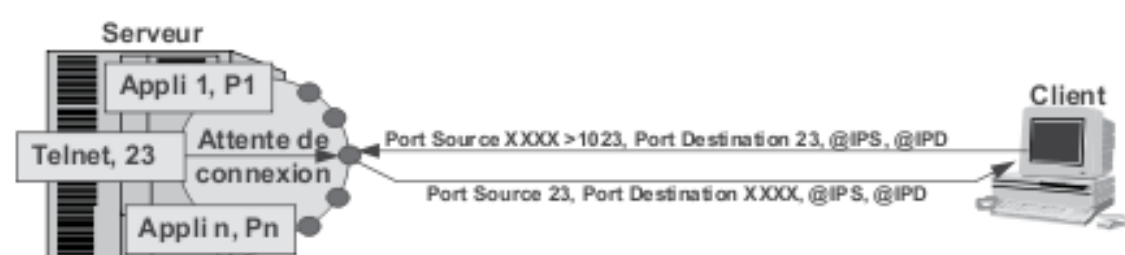


Figure 14.4 Le référencement de la connexion de transport.

14.2 Etablissement de la connexion de transport

TCP définit deux modes d'ouverture : le mode passif et le mode actif. Dans le mode passif, TCP est en attente d'une demande d'ouverture en provenance d'un autre système (défini ou non). Dans le mode actif, TCP adresse une demande de connexion à un autre système identifié (figure 14.5).

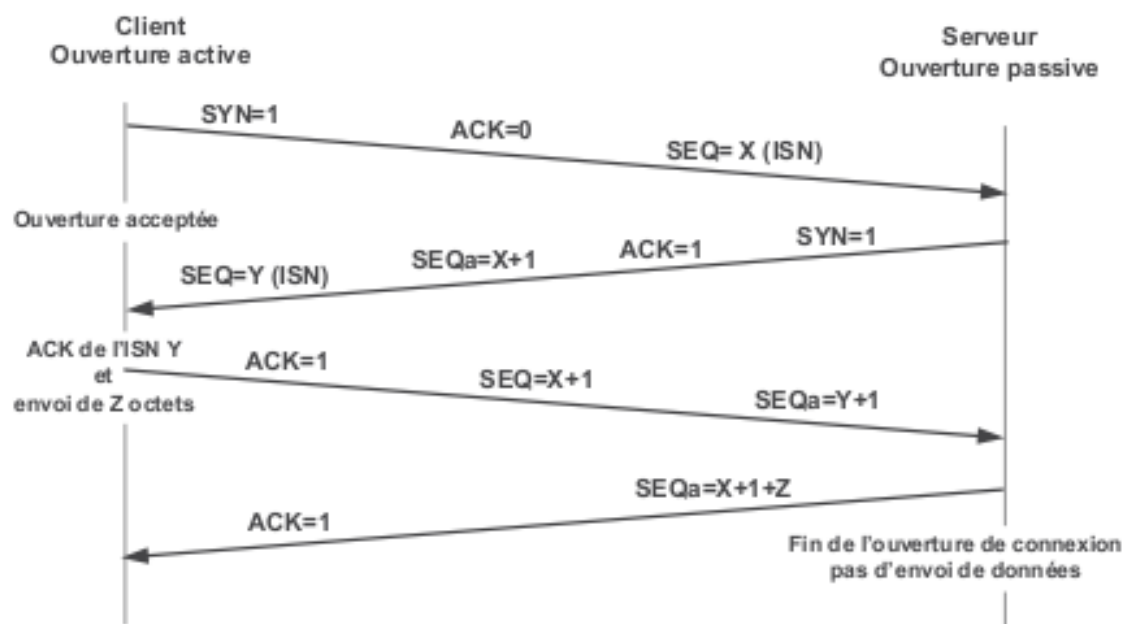


Figure 14.5 L'ouverture d'une connexion en trois temps.

L'ouverture et la fermeture de connexion sont gérées par 2 bits de l'en-tête TCP : le bit SYN et le bit FIN. À l'instar d'HDLC, TCP numérote les données transmises. Si HDLC dénombre des trames, TCP compte des octets ainsi, pour chaque segment TCP transmis, il identifie la position (*Offset*), dans le flux global de données du premier octet, décomptée à partir du premier octet du premier segment transmis. Contrairement à HDLC, les compteurs ne sont pas initialisés à zéro en début de transmission, mais à une valeur dite ISN (*ISN, Initial Sequence Number*). Lors de l'ouverture de connexion, chaque entité informe son correspondant de sa valeur locale (initialisation des compteurs de part et d'autre).

La connexion de transport est *full duplex* mais chaque sens est indépendant l'un de l'autre. Il peut être mis fin à un échange (figure 14.6) dans un sens

sans pour cela mettre fin à l'échange dans l'autre sens (semi-fermeture). La fin d'un échange dans un sens peut donc être tout à fait dissociée de la fermeture. La fermeture définitive n'est effective que lorsque chaque sens a clos son échange. La fermeture est dite douce ou négociée évitant ainsi toute perte de données dans le réseau (il n'y a pas de couche session pour contrôler le dialogue).

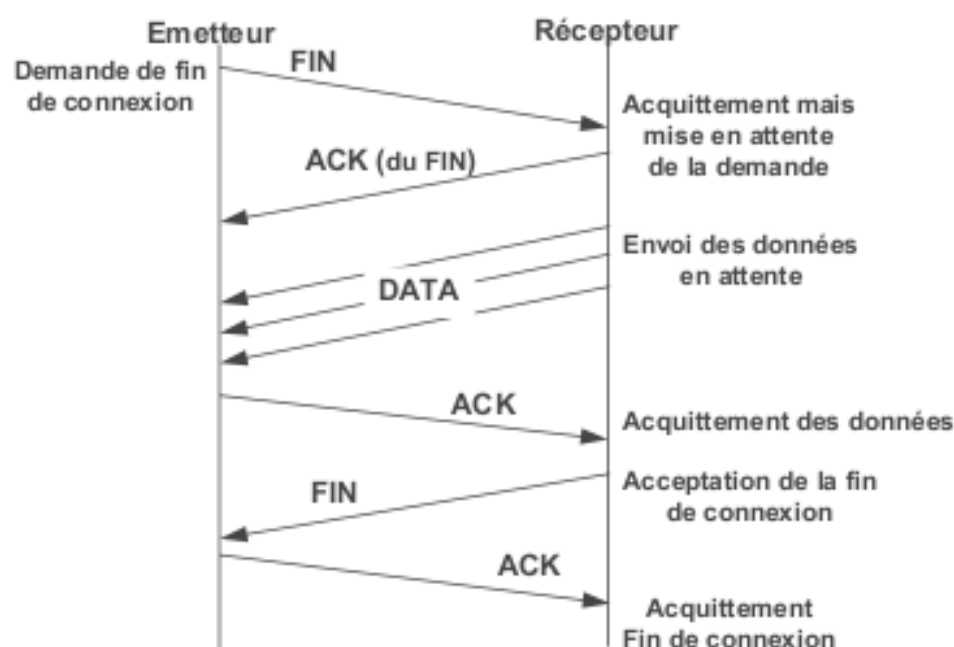


Figure 14.6 La fermeture négociée de TCP.

14.3 Le mécanisme contrôle de l'échange

14.3.1 L'ISN

À l'instar d'autres protocoles, pour fiabiliser la connexion et détecter d'éventuels segments perdus, TCP numérote les octets transmis (numéros de séquence sur 32 bits) ; ainsi, il identifie, dans chaque segment transmis, la position (*Offset*) dans le flux de données du premier octet de ce segment.

TCP attribue à chaque connexion un numéro de séquence initial (**ISN**, *Initial Sequence Number*) différent, le décomptage des octets transmis ne

démarre pas à 0, mais à une valeur initialisée à partir de l'horloge interne de la machine.

14.3.2 La gestion des acquittements

TCP utilise le mécanisme de l'anticipation et met en œuvre la technique de la fenêtre glissante, chaque bloc émis doit être acquitté immédiatement. Cependant, supposons une application de type terminal asynchrone (Telnet). L'application cliente émet, pour chaque touche frappée, un segment contenant le caractère et l'application serveur renvoie ce caractère en écho (écho distant), la couche transport envoie un accusé de réception pour chaque message émis (figure 14.7 gauche).

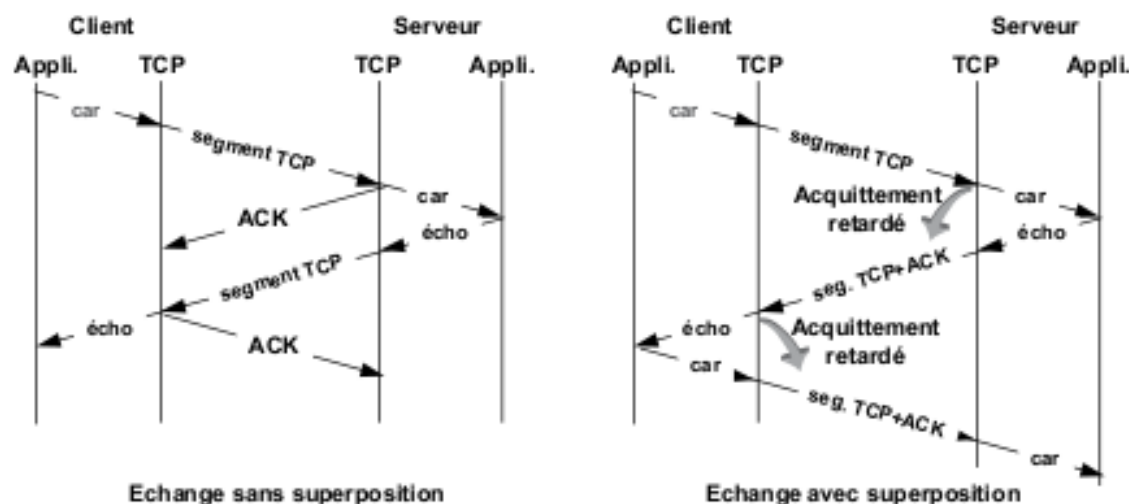


Figure 14.7 Principe de la superposition.

Si on considère l'ensemble des en-têtes TCP et IP, soit 40 octets d'en-tête par message, pour 1 caractère d'utile on a transmis 122 caractères ! Un moyen simple d'augmenter l'efficacité de cette transmission serait de profiter de l'envoi du caractère en écho pour transmettre l'information d'acquiescement (figure 14.7 droite). C'est le principe de la superposition (*Piggybacking*), un même segment peut transporter non seulement des données, mais aussi l'acquiescement de segments précédents reçus.

Afin d'optimiser ce mécanisme, TCP retarde les acquittements dans l'attente de données à transmettre (*delayed ACK*). Pour ne pas nuire à

l'interactivité des applications, ce délai d'attente est généralement compris entre 200 ms et 500 ms (valeur implicite 200 ms). Le comportement de TCP est décrit dans le tableau 14.2.

Tableau 14.2 Les mécanismes de reprise de TCP.

Type de réception	Comportement de TCP
Réception d'un segment attendu (en séquence et sans trou), les précédents ayant été acquittés.	Mise en attente de 200 à 500 ms (algorithme de Nagle) puis envoyer un éventuel segment en attente (piggybacking), à défaut acquitter le segment.
Réception d'un segment attendu pendant le délai de garde (1 ACK en attente).	Envoyer un ACK cumulant l'acquittement du segment précédent et de celui qui vient d'être reçu.
Réception d'un segment avec un numéro de séquence supérieur au numéro attendu (déséquencement ou perte de segments).	Renvoyer le dernier ACK portant le numéro de séquence attendu (Acquittement dupliqué).
Réception d'un segment dans l'ordre qui remplit partiellement ou totalement un trou (réception précédente supérieure au N° de séquence attendu).	Si le trou est rempli totalement, envoyer immédiatement l'acquittement, sinon mise en attente.

TCP n'acquitte que les données correctement reçues. La retransmission des données erronées s'effectue sur temporisation. La figure 14.8 illustre ce fonctionnement. Pour simplifier la compréhension, les données ont été numérotées en N° de segment et non en octets. Dans cet exemple, le destinataire a correctement reçu les segments 1 et 2, le segment 3 erroné est ignoré. Compte tenu du mécanisme d'anticipation, le destinataire reçoit les segments 4, 5 et 6. La réception du segment 4, hors séquence, provoque l'émission immédiate d'un acquittement précisant le numéro de segment attendu, c'est-à-dire le 3. Toutes les données reçues hors séquencement sont mémorisées, mais non expressément acquittées, c'est toujours le dernier segment reçu en séquence qui est acquitté. Ainsi l'émetteur reçoit plusieurs acquittements du segment 2 (segment attendu 3). L'émetteur ne réagit pas immédiatement, ce n'est qu'après la réception de trois accusés de réception d'un même segment avant l'échéance de la temporisation (RTO), qu'il effectue la retransmission du segment demandé (N° d'octets).

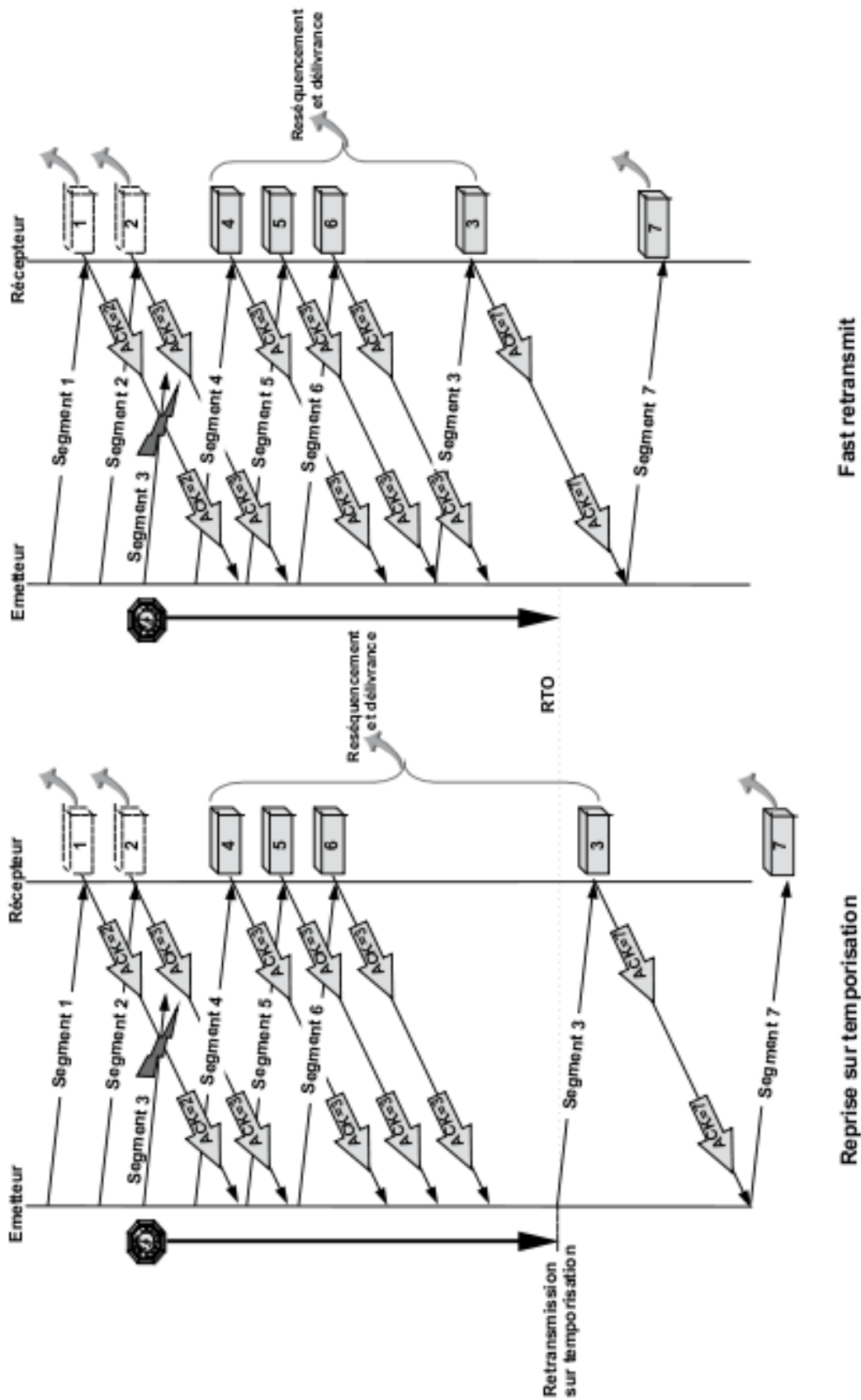


Figure 14.8 La reprise sur temporisation.

Dans notre cas, ne recevant pas d'acquittement du segment 3 (N° attendu 4) avant l'échéance du temporisateur, l'émetteur renvoie le segment 3, et cesse toute émission en attente de l'acquittement de ce bloc. À réception du segment 3, le récepteur acquitte l'ensemble des segments correctement reçus (acquittement cumulatif). Ainsi, l'acquittement indique que le destinataire est en attente du segment 7. Cette technique dénommée « **Fast retransmit** » permet d'optimiser la transmission en évitant une retransmission de données correctement reçues.

En cas de reprise, le TCP émetteur peut émettre un segment plus grand que celui devant être retransmis. De ce fait, certains octets des segments suivants (5, 6, 7) déjà reçus pourront éventuellement l'être de nouveau. C'est le TCP destination qui gèrera l'éventuel recouvrement des données.

14.3.3 La gestion des temporisations

Les délais d'acheminement dans un réseau IP dépendent de nombreux facteurs et en particulier :

- ▶ du débit du ou des réseaux traversés (figure 14.9),
- ▶ de la charge de ceux-ci,
- ▶ de la stabilité de route (mode datagramme)...

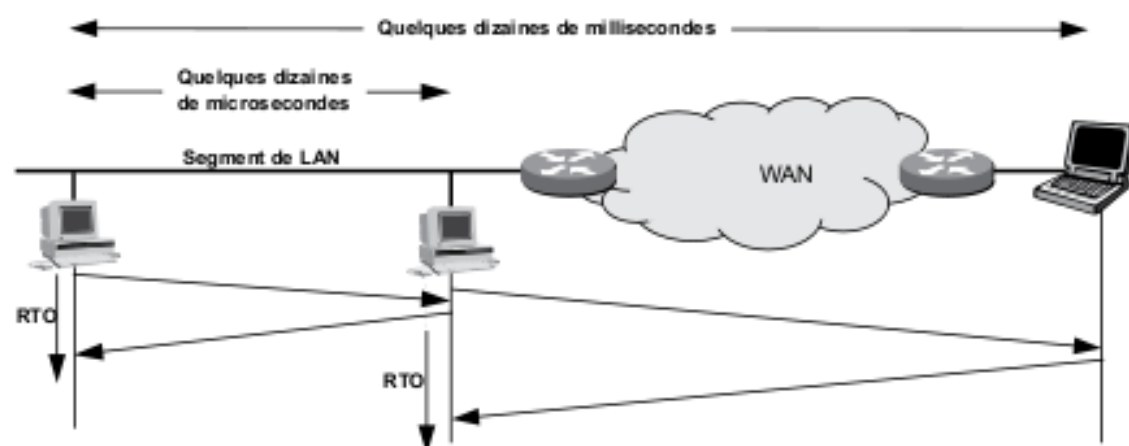


Figure 14.9 TCP et la difficulté de détermination du RTO.

Dans ces conditions, il ne peut être fait d'hypothèses précises sur les délais d'acheminement et par conséquent sur la valeur des temporisateurs. La

reprise sur temporisation (**RTO**, *Retransmission Time Out*) est l'un des facteurs de performance d'un protocole, un délai trop faible conduit à des retransmissions inutiles, un délai trop important à une baisse des performances globales. TCP devant s'adapter à tout type de réseau met en œuvre un mécanisme dynamique de détermination et d'ajustement du RTO. Ce mécanisme permet la prise en compte des variations du temps de transfert dans le réseau, variations pouvant être dues à une variation de charge dans le réseau voire un changement de route.

14.4 L'option d'estampille horaire

Les performances globales de TCP reposent sur la détermination du RTT. À cet effet, TCP réalise cette mesure en calculant le temps entre l'émission d'un segment et la réception de son acquittement.

L'option d'estampille horaire (*timestamps*) permet à TCP de déterminer avec précision le RTT du réseau, temps de traitement du destinataire compris. La figure 14.10 représente le format de l'option d'estampille horaire.

Code=8 1 octet	Longueur=10 1 octet	Estampille horaire 4 octets	Echo estampille horaire 4 octets
-------------------	------------------------	--------------------------------	-------------------------------------

Figure 14.10 Le format de l'option estampille horaire.

Le TCP émetteur indique sur 32 bits l'instant d'émission dans le champ Estampille horaire. Le TCP destinataire retourne en écho son instant d'émission. Le RTT est déterminé par l'émetteur par différence entre l'instant d'émission de l'estampille horaire et l'instant de réception de son écho.

15

Le segment TCP et les mécanismes associés

15.1 La structure du segment TCP

TCP ne définit qu'un seul format de segment, l'en-tête de TCP (figure 15.1) contient toutes les informations nécessaires à la gestion de la connexion de transport (commandes), le transfert de données, la gestion des acquittements, le contrôle de flux et de congestion.

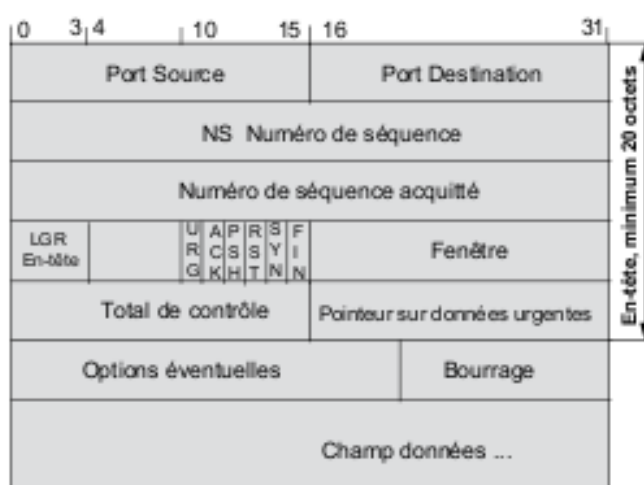


Figure 15.1 La structure de la TPDU ou segment TCP.

Les différents champs de l'en-tête du segment sont :

- **Numéros de port** (deux fois 16 bits), valeurs spécifiées à la connexion et identifiant celle-ci. Le port du programme appelé est soit connu, soit défini lors d'une phase d'identification (login) et passé à l'appelant en réponse.
- **Numéro de séquence** sur 32 bits, initialisé à une valeur aléatoire (ISN, *Initial Sequence Number*). Le numéro de séquence indique le rang du premier octet du segment transmis.

- ▶ Le **numéro de séquence acquitté** indique le numéro du prochain octet attendu.
- ▶ **Longueur de l'en-tête** ou *offset* (4 bits) indique la longueur de l'en-tête du segment TCP en multiples de mots de 32 bits (4 octets).
- ▶ Le champ **Drapeau** contient six indicateurs :
 - **URG** valide le contenu du champ « pointeur sur données urgentes ». Si ce bit est positionné à 1, le segment transporte des données à traiter en priorité ;
 - **ACK**, à 1 il valide le contenu du champ « numéro de séquence acquitté » ;
 - **PSH**, l'indicateur *Push* permet de demander à l'émetteur de transmettre immédiatement les données en attente et au destinataire de les traiter sans retard ;
 - **RST**, demande au destinataire de réinitialiser la connexion ;
 - **SYN**, à 1 ce bit correspond à une demande de connexion ;
 - **FIN**, ce drapeau correspond à une demande de déconnexion émise par l'un des interlocuteurs.
- ▶ Le champ **fenêtre** (2 octets) indique, en octet, la valeur de la fenêtre en réception (crédit alloué) ;
- ▶ Le **total de contrôle** (*Checksum*) est calculé sur l'ensemble du segment TCP, en-tête compris (voir § suivant) ;
- ▶ Le **pointeur sur données urgentes** pointe sur le dernier octet urgent du champ de données, les données urgentes sont traitées, en priorité, par le destinataire.

15.2 Le contrôle d'erreur

TCP et UDP utilisent la même technique de contrôle d'erreur. Pour calculer le total de contrôle, TCP (émetteur et récepteur) adjoint au segment TCP (ou UDP) un pseudo-en-tête IP contenant, notamment, les adresses IP source et destination. Le total de contrôle porte donc sur les données, l'en-tête TCP (ou UDP) et le pseudo-en-tête IP (figure 15.2).

Cette manière de procéder garantit l'intégrité des données et la délivrance au bon destinataire, mais est en violation avec la règle d'indépendance des couches.

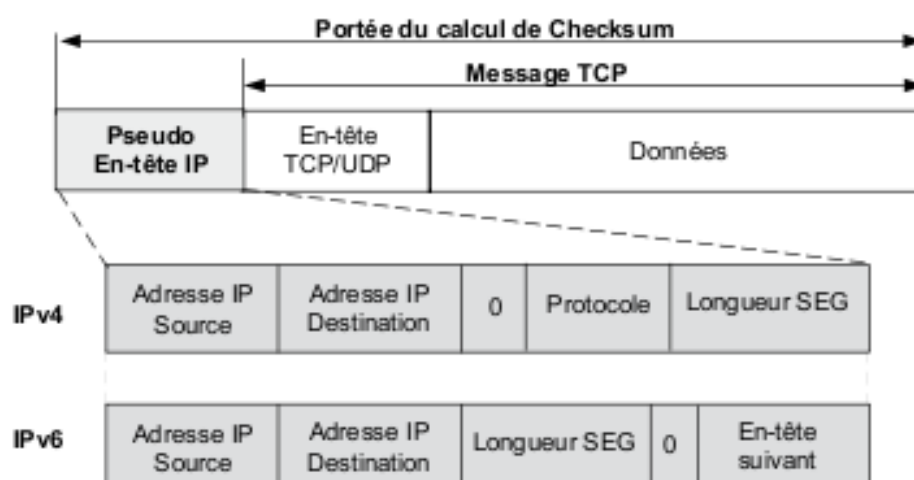


Figure 15.2 La portée du contrôle d'erreur dans UDP et TCP.

Le mode de calcul est identique à celui mis en œuvre par IP. Il s'agit du complément à 1 de la somme en complément à 1 des mots de 16 bits de l'intégralité du segment TCP (UDP), pseudo en-tête inclus. Le pseudo-en-tête n'est pas transmis. Notons que les en-têtes IPv4 et IPv6 étant différentes, le TCP au-dessus de ces protocoles réseaux l'est. Les piles IPv4 et IPv6 cohabitent et sont indépendantes.

15.3 La taille des segments

Lors de l'établissement de la connexion de transport, l'option **MSS** (*Maximum Segment Size*) permet l'annonce de la taille maximale de segment que le système d'extrémité peut admettre. À défaut d'annonce si l'application est locale, la MSS correspondra à la MTU de l'interface réseau diminuée des en-têtes TCP et IP. La MTU minimale (figure 15.3) est fixée, dans IPv4, à 576 octets (1 280 IPv6).

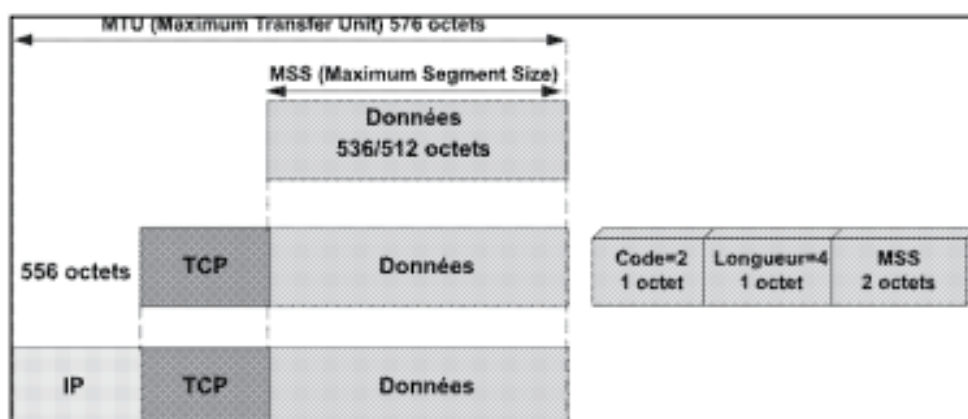


Figure 15.3 La relation entre MTU et MSS, et le format de l'option.

15.4 Le TCP et les réseaux à haut débit

Les performances maximales d'un protocole à anticipation sont atteintes lorsque le crédit de l'émetteur (fenêtre), exprimé en temps d'émission, correspond au (RTT) :

$$\text{Fenêtre en bits} = \text{débit du lien (réseau)} \times \text{RTT}$$

Le champ taille de la fenêtre est limité à 16 bits soit une capacité de 65 535 octets. Dans ces conditions, TCP est mal adapté aux réseaux dont le produit Débit par RTT est important (réseaux dits *Long Fat Network*, LFN, ou réseau éléphant !). Une fenêtre plus petite que la capacité du réseau provoque, dans l'attente d'un acquittement, l'arrêt des émissions. Pour adapter TCP aux réseaux à haut débit où la latence est importante (liaison satellite), la capacité de numérotation de la fenêtre a été accrue. Invoquée lors de la connexion, l'option d'échelle de fenêtre (RFC 1323) indique en puissance de 2 le décalage de la fenêtre. Ainsi, si la valeur est de 1, la fenêtre utilisée sera de $65\,535 \times 2^1$, si le décalage est de 2, la valeur de fenêtre est de $65\,535 \times 2^2$... La figure 15.4 illustre le format de l'option d'échelle de fenêtre.

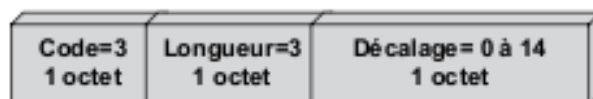


Figure 15.4 L'option échelle de fenêtre.

16 TCP/UDP et le multimédia, le contrôle de flux et de congestion

16.1 Définitions

Les contrôles de flux et de congestion sont des mécanismes similaires souvent confondus, aussi rappelons que :

- ▶ le **contrôle de congestion** consiste à contrôler le débit d'émission d'une source en fonction de la capacité du sous-réseau physique réel ;
- ▶ le **contrôle de flux** consiste à contrôler le débit d'émission d'une source en fonction de la capacité de réception du destinataire.

Dans un réseau en mode datagramme, ces contrôles ne peuvent être réalisés dans le réseau, ces mécanismes sont alors reportés sur les machines d'extrémité (algorithme distribué). À cet effet, TCP gère deux fenêtres, une de contrôle de congestion (**cwnd**), déterminée par l'émetteur, et une de contrôle de flux (**rwnd**), crédit octroyé par la source. Ces deux fenêtres sont indépendantes l'une de l'autre, le TCP émetteur limite ses émissions à la plus petite des deux valeurs :

$$W_{\text{émission}} = \text{Min} [\text{rwnd (receive window, Indication du distant)}, \\ \text{cwnd (congestion window, gérée localement)}]$$

Hors état de congestion, la fenêtre de contrôle de flux et celle de contrôle de congestion sont identiques.

16.2 Le contrôle de flux

16.2.1 Principe

Le contrôle de flux est réalisé par un mécanisme à **fenêtre dynamique** ou **contrôle de flux explicite**. Le crédit accordé est un crédit d'octets (*rwnd*).

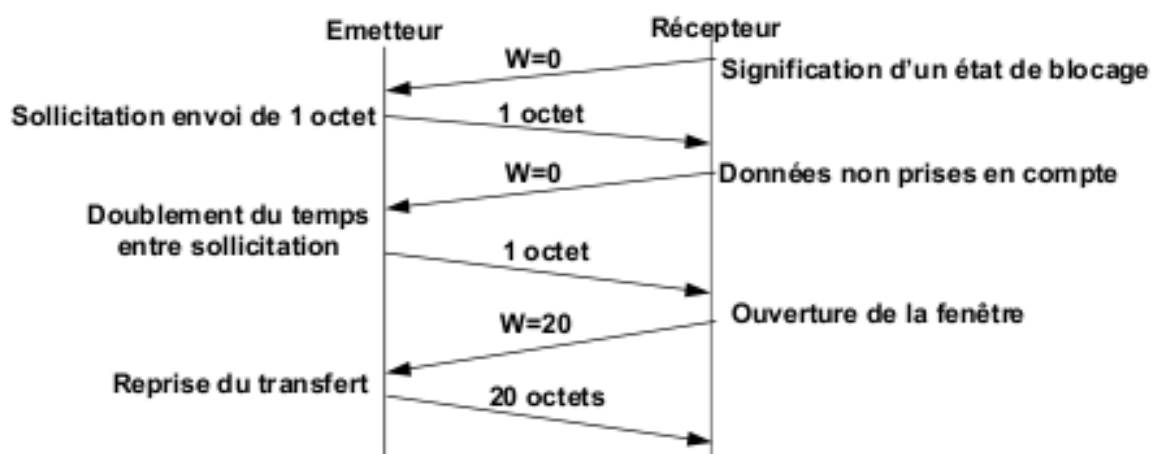


Figure 16.1 Principe du contrôle de flux par TCP.

Dans la figure 16.1, le *buffer* plein, le récepteur en informe l'émetteur en positionnant dans l'acquittement le champ Fenêtre à 0. L'émetteur cesse ses émissions. Pour éviter la fermeture de la connexion sur inactivité, l'émetteur sollicite périodiquement le récepteur en lui envoyant 1 octet de données (non pris en compte). Le processus se poursuit jusqu'à ce que le récepteur débloque la situation en ouvrant la fenêtre.

16.3 Le contrôle de la congestion

16.3.1 Principe

Sur un réseau en mode datagramme, la congestion ne peut être détectée que par le récepteur suite à un allongement des temps de transit (acquittements retardés) ou à la perte de données (les acquittements n'arrivent pas). TCP considère que toute perte de segment est due à un état de congestion.

16.3.2 Le contrôle de congestion de bout en bout

À chaque segment perdu, le TCP émetteur réduit ses émissions par réduction dichotomique de la fenêtre de congestion (*cwnd*) et, pour ne pas surcharger le réseau de retransmissions peut-être inutiles, il augmente la valeur du *timer* de retransmission (RTO). Après congestion, pour éviter un retour à l'état de congestion, TCP adopte une technique de redémarrage simple, la fenêtre de congestion est fixée à 1 et s'accroît à chaque réception d'acquittement (démarrage lent).

■ Le démarrage lent et l'évitement de congestion

Plusieurs techniques de démarrage et de traitement de la congestion peuvent être mises en œuvre par TCP/IP. Une phase de démarrage (figure 16.2), dite démarrage lent (**Slow start**), fixe la fenêtre de congestion à un segment (*cwnd* = 1) et détermine une valeur maximale de la fenêtre de congestion **SSthresh** (*Slow Start Threshold*) fixée à 65536. À chaque accusé de réception reçu, la fenêtre de congestion est doublée (croissance exponentielle). Lors de la perte d'un segment, TCP considère avoir atteint la limite de la capacité disponible sur le lien, fixe une nouvelle valeur **SSthresh** ($SSthresh = cwnd/2$) et entreprend une nouvelle phase de démarrage lent. Pour éviter un retour rapide à l'état de congestion, dès que les segments émis ont atteint la valeur contenue dans la variable **SSthresh**, la progression n'est plus exponentielle mais linéaire (évitement de congestion ou **congestion avoidance**). Cet algorithme (TCP Tahoe) utilise mal les ressources du réseau puisqu'il redémarre une phase d'émission à 1 segment. Une autre version dite « Reno », la plus implémentée, adopte un comportement différent selon le mode de détection d'un paquet perdu :

- ▶ Si la perte est détectée par la réception de 3 Ack dupliqués (*Fast retransmit*, RFC 2581), TCP Reno ne refait que la phase d'évitement de congestion (*Congestion avoidance*), cette technique (**Fast recovery**) autorise un retour plus rapide aux valeurs optimales ;

- Si la perte est détectée par échéance du RTO, TCP Reno exécute la phase de *slow start*.

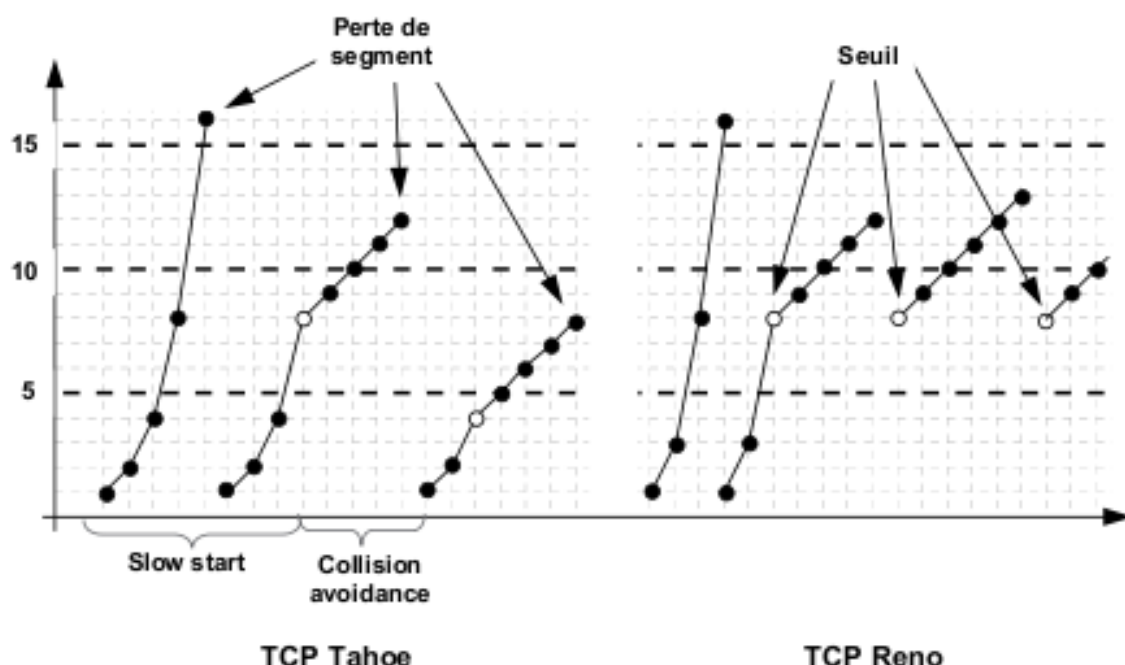


Figure 16.2 Le démarrage lent et l'évitement de congestion.

■ Performances de TCP

En régime stationnaire (TCP Reno), **cwnd** oscille autour de la fenêtre optimale (figure 16.2). Dans ces conditions, on montre que le débit maximal d'une liaison TCP dépend de la taille maximale en octets du segment de données (MSS), du délai d'acheminement aller et retour (RTT) mais aussi du taux de perte de segments (p) :

$$\text{Débit}_{(\text{octets par seconde})} = \frac{1,22 \cdot \text{MSS}}{\text{RTT} \cdot \sqrt{p}}$$

Interprétation : le mécanisme de redémarrage, lors de la perte d'un paquet quelle qu'en soit la cause pénalise fortement les performances, TCP est donc mal adapté aux réseaux peu fiables et notamment aux réseaux sans fil !

16.3.3 Le contrôle de congestion par information de l'état du réseau

L'état de congestion se manifeste dans un routeur par le débordement des files d'attente, tout nouveau paquet arrivant est alors éliminé. Tous les flux transitant par ce routeur sont pénalisés et vont réduire leur fenêtre de congestion et, selon leur version de TCP, pratiquer un redémarrage lent, un *Fast recovery*... Pour éviter cette synchronisation globale, il convient d'anticiper l'état de congestion en demandant, aux seules sources les plus contributives, de réduire leur émission, c'est le principe du RED (*Random Early Detection*).

16

■ Mécanisme du RED (RFC 2309)

Le mécanisme du RED consiste à mesurer en permanence la longueur moyenne des files d'attente (L_{moy}) et, en fonction de celle-ci éliminer, selon une probabilité variant avec la taille de la file, le paquet entrant (figure 16.3).

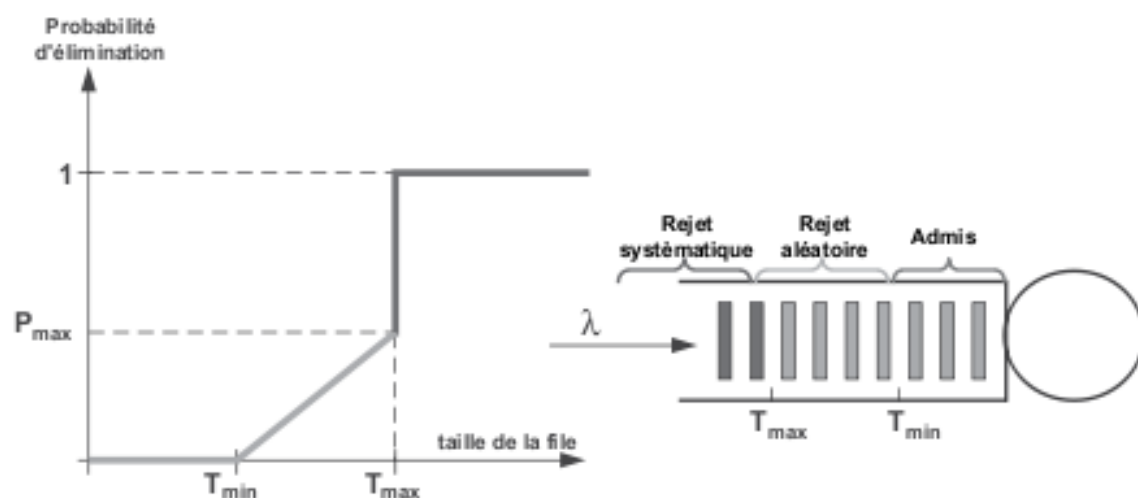


Figure 16.3 Probabilité de rejet en fonction de la taille de la file d'attente.

Le RED pénalise d'autant plus les flux que ceux-ci sont importants, évitant ainsi une synchronisation globale du réseau.

■ TCP ECN (Explicit Congestion Notification)

Afin d'éviter les éliminations préventives du RED, TCP ECN (RFC 3168) implémente un mécanisme d'alerte provoquant une réduction des fenêtres de congestion avant que celle-ci n'apparaisse ou que le RED n'entre en service, c'est un mécanisme d'anticipation. TCP ECN substitue une alerte à une élimination.

La difficulté de réaliser un mécanisme d'anticipation sur un réseau datagramme provient du fait que la congestion se produit au niveau 3 (niveau réseau) et que la détection se fonde sur la perte de données au niveau transport (niveau 4). TCP ECN utilise une version modifiée d'IP et de TCP qui viole l'indépendance des couches, l'alerte est transmise au niveau 3 (bits ECN) et gérée au niveau 4 (bits CWR et ECE). La figure 16.4 représente le format des unités de données de TCP ECN.

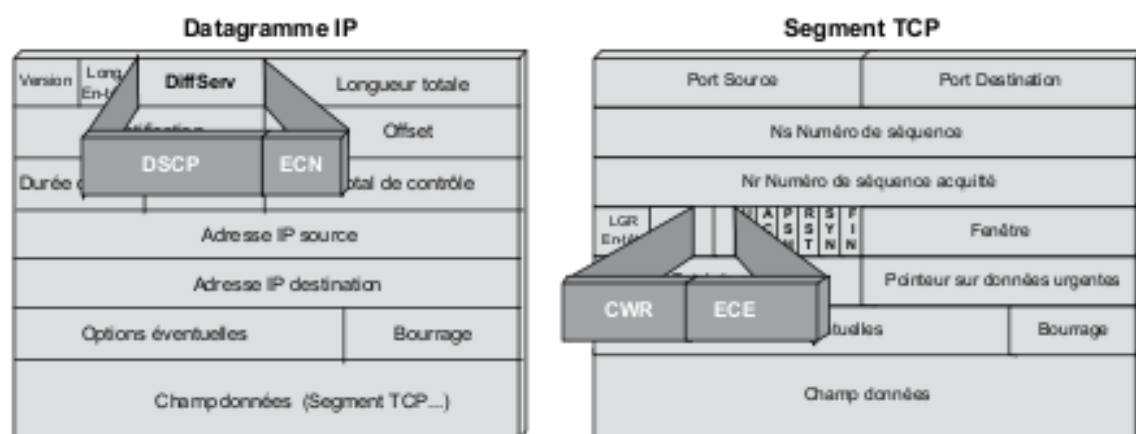


Figure 16.4 Les bits de gestion TCP/IP ECN.

La figure 16.5 illustre le scénario ECN complet. Lors de l'ouverture des connexions TCP, chaque correspondant informe son distant de sa capacité à gérer l'ECN (CWR = 1, ECE = 1). Le distant acquitte cette information au niveau réseau (ECT(1) = 10). En l'absence de congestion les bits CWR et ECE sont positionnés à zéro, tandis que les bits ECN des datagrammes indiquent la capacité à gérer l'ECN (ECT = 10). Lors de la traversée d'un routeur en état de pré-congestion, les bits ECN du datagramme IP sont positionnés à 11 (ECT = 11 ou état CE).

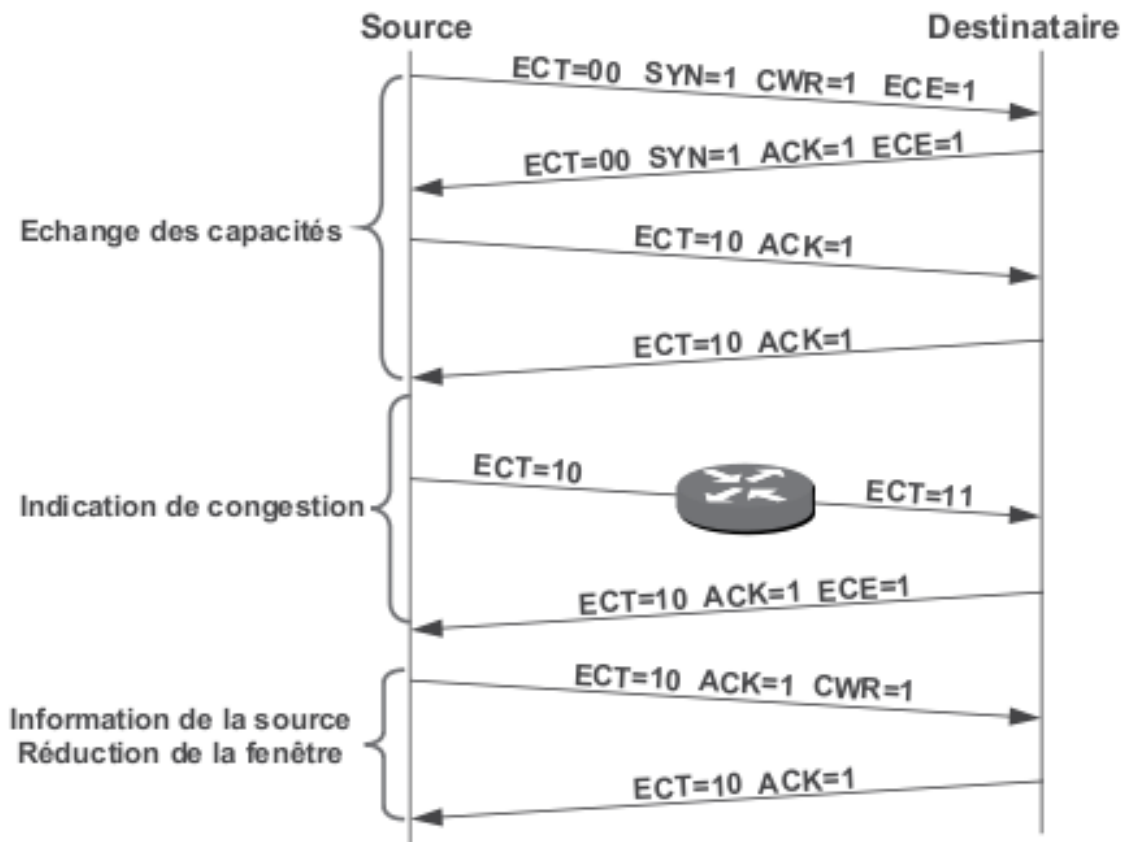


Figure 16.5 Schématisation du dialogue complet d'un échange ECN.

Le destinataire informe la source que son flux traverse un nœud en état de pré-congestion ($ECE = 1$) ; celle-ci réduit sa fenêtre (*Fast retransmit*) et en informe son distant ($CWR = 1$). Dans les messages suivants, le destinataire ne redemande pas, dans l'intervalle de temps d'un RTT, de nouvelle réduction de la fenêtre.

16.4 UDP dans IPv4

Certaines applications et en particulier les applications « temps réel » nécessitent des temps de traitement optimisés, d'autres n'ont que très peu de données à envoyer, comme le service de noms (DNS), d'autres encore n'ont nullement besoin d'un service sécurisé comme les informations de gestion des réseaux (SNMP). Aussi, un mode de transport allégé a-t-il été défini : UDP (*User Datagram Protocol*, RFC 768).

Le segment (datagramme) UDP ne contient que les champs ports source et destination, longueur totale du datagramme (en-tête compris sur 2 octets), total de contrôle (2 octets) et enfin les données utilisateurs (figure 16.6). Le *checksum* UDP est calculé comme celui de TCP pseudo-en-tête IP compris. Son utilisation est facultative, en cas de non-utilisation, le champ *Checksum* est mis à zéro.

Port Source UDP	Port destination UDP
Longueur Segment	Checksum UDP
Données	

Figure 16.6 Le format du segment UDP.

16.5 UDP dans IPv6

L'une des simplifications majeures d'IPv6 concerne l'abandon du calcul du checksum sur l'en-tête IP. TCP inclut dans son calcul de checksum un pseudo-en-tête IP protégeant ainsi les adresses IP d'éventuelles erreurs. En conséquence dans IPv6, pour assurer une protection minimale aux adresses IP source et destination, le calcul du total de contrôle dans UDP a été rendu obligatoire.

16.6 Conclusion

La pile protocolaire TCP/IP offre deux modes de transport, un mode assuré, TCP, qui garantit la délivrance de données, assure le contrôle de flux et de congestion, et un mode allégé, UDP, gage de performance mais sans garantie de délivrance. Les applications traditionnelles utilisent TCP pour sa fiabilité, tandis que les applications dites temps réel comme la voix et la vidéo s'appuient sur UDP. Dans un réseau à faible bande passante, les applications de type vidéo, non soumises au contrôle de flux, pénalisent les applications traditionnelles (figure 16.7).

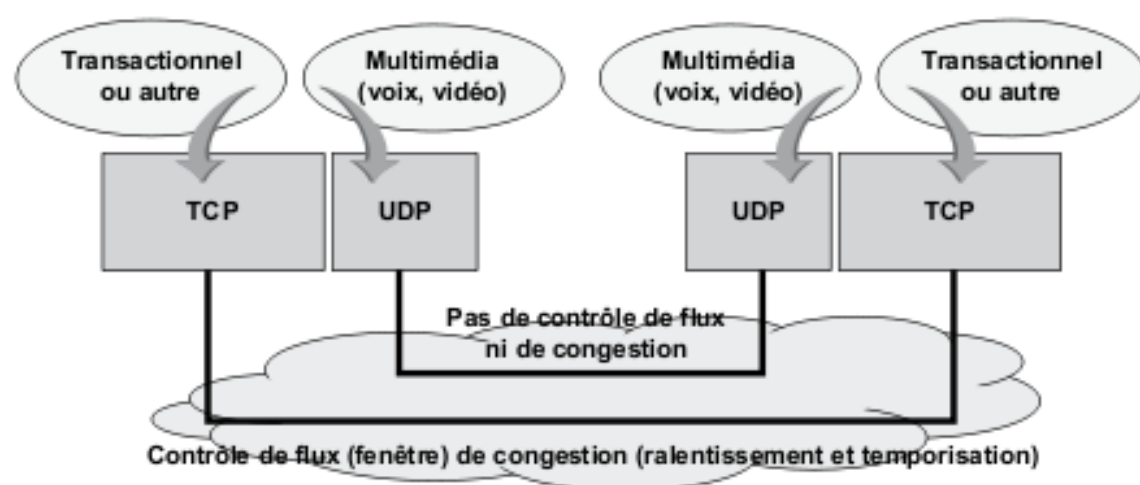


Figure 16.7 Le partage de la bande passante par TCP et UDP.

6

TCP/IP

**utilitaires
et applications**

17

Les utilitaires de la couche réseau

17.1 Le protocole ICMP

17.1.1 Généralités (ICMPv4)

Le protocole ICMP (*Internet Control Message Protocol*, RFC 792) permet d'informer la source d'une erreur réseau (message d'erreur) ou de formuler une demande d'état à un système (message d'information). Les messages ICMP sont encapsulés dans un datagramme IP (Protocole = 1). La figure 17.1 représente la structure du message d'erreur ICMP et fournit quelques exemples de codage des différents champs.

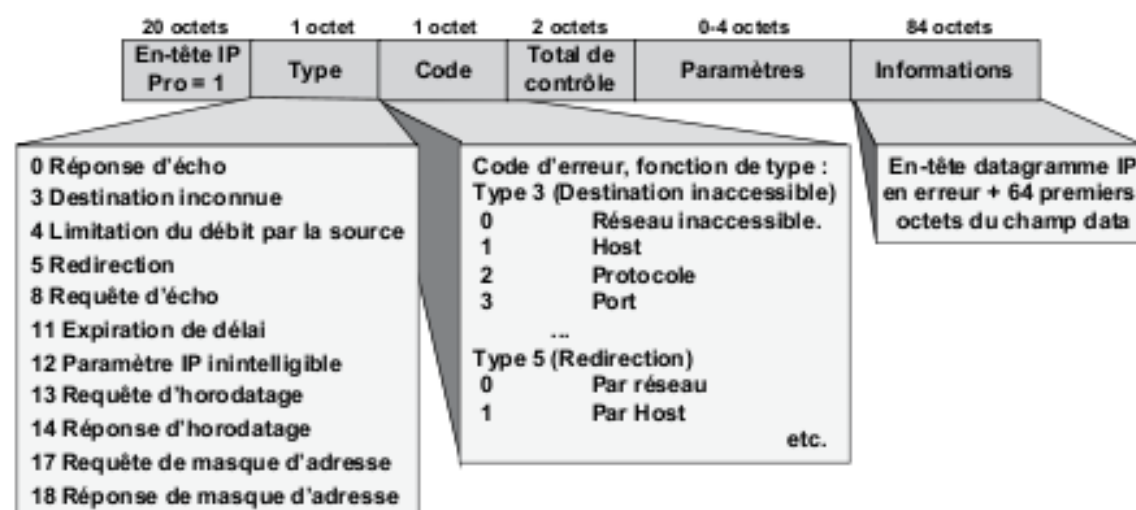


Figure 17.1 La structure du message d'erreur d'ICMP.

Le protocole ICMP ne fiabilise pas un réseau IP, c'est un protocole d'information. Ses différentes fonctions sont aussi utilisées par certains utilitaires. Par exemple, la commande PING (voir § 1.1.3 ci-dessous) uti-

lise les messages de demande de réponse d'écho (type 0 et 8). Le protocole **NTP** (*Network Time Protocol*) utilise la commande *ICMP Timestamp* (type 13 et 14) pour synchroniser les horloges du réseau.

17.1.1.2 ICMPv6

À l'instar d'ICMPv4, ICMPv6 informe la source d'une erreur réseau (message d'erreur, type < 127), permet de formuler une demande d'état à un système (message d'information, type > 127) ou enfin d'offrir de nouvelles fonctionnalités d'autoconfiguration (découverte des voisins, gestion des groupes *multicast*...). Les messages ICMP sont encapsulés dans un datagramme IP (Protocole = 58). La figure 17.2 représente la structure du message ICMPv6 et fournit quelques exemples de codage des différents champs.

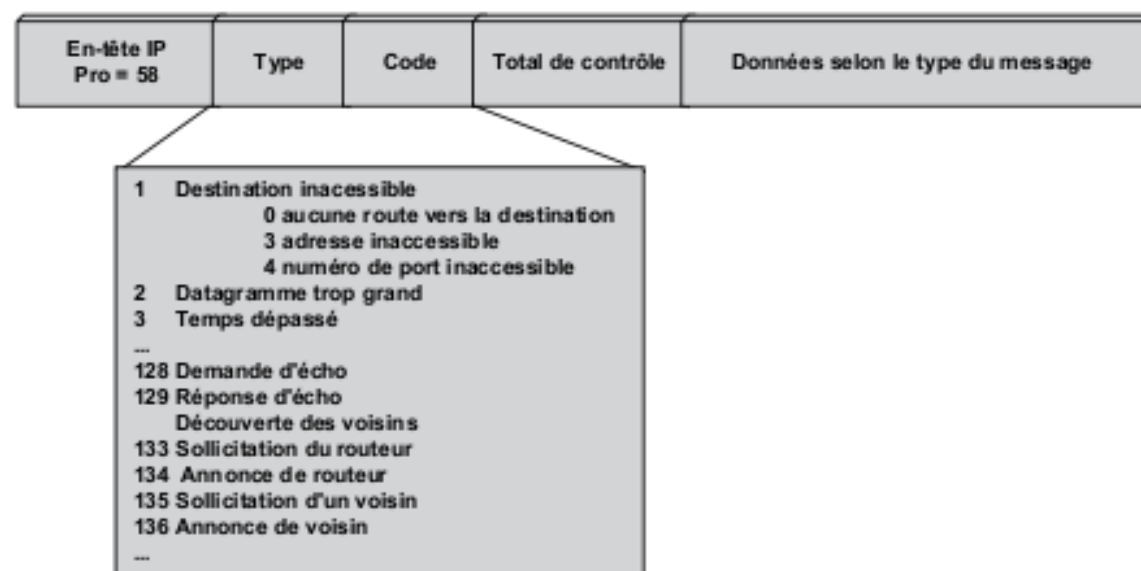


Figure 17.2 Message ICMPv6.

17.1.1.3 Utilitaires utilisant les messages ICMPv4

■ L'utilitaire PING

L'utilitaire **PING** (*Packet INternet Groper*) permet de tester l'accessibilité d'un système. PING envoie un message ICMP de demande d'écho vers la machine cible, celle-ci retourne une réponse d'écho (figure 17.3).



Figure 17.3 Principe de la commande PING.

La commande PING permet de tester une liaison TCP/IP de bout en bout :

- ▶ PING sur l'adresse 127.0.0.1 teste l'installation de la pile TCP/IP sur la machine source ;
- ▶ Sur l'adresse IP de la machine source, elle vérifie que cette station est correctement configurée ;
- ▶ Sur l'adresse de la passerelle par défaut, elle contrôle la validité du masque de sous-réseau et la configuration de la passerelle par défaut ;
- ▶ PING sur l'adresse de l'interface de sortie (LS locale) valide la configuration de cette interface ;
- ▶ Sur l'adresse de LS distante, elle s'assure que le lien WAN est établi et que les routeurs local et distant sont correctement configurés vis-à-vis du réseau source ;
- ▶ Sur l'adresse de station distante, elle valide la configuration de bout en bout.

La commande PING émet une série de demandes d'écho, la figure 17.4 illustre la réponse à une commande PING sur une cible d'adresse 195.221.126.186. Chaque demande d'écho a été émise avec un champ données de 32 caractères. L'information temps indique la durée entre l'émission de la requête et la réception de la réponse, c'est une approximation du RTT. Le champ TTL, non présent dans toutes les implémentations, reproduit la valeur du champ TTL du datagramme IP reçu. Cette valeur est toujours délicate à interpréter, elle dépend de la valeur initiale du TTL fixée par la machine cible. Certaines implémentations fixent à 255 le TTL des réponses d'Echo d'ICMP.

C:> Ping 195.221.126.186

Envoi d'une requête 'ping' sur 195.221.126.186 avec 32 octets de données :

Réponse de 195.221.126.186 : octets=32 temps=135 ms TTL=53

Réponse de 195.221.126.186 : octets=32 temps=140 ms TTL=53

Réponse de 195.221.126.186 : octets=32 temps=156 ms TTL=53

Réponse de 195.221.126.186 : octets=32 temps=156 ms TTL=53

Statistiques Ping pour 195.221.126.186:

Paquets : envoyés = 4, reçus = 4, perdus = 0 (perte 0%),

Durée approximative des boucles en millisecondes :

Minimum = 135 ms, Maximum = 156 ms, Moyenne = 146 ms

Figure 17.4 Exemple de commande PING.

■ L'utilitaire Traceroute (Tracert)

L'utilitaire Traceroute dû à Van Jacobson permet de découvrir la route empruntée par un datagramme entre une source et une machine cible. L'utilitaire Traceroute émet trois datagrammes IP successifs avec un TTL incrémenté d'un à chaque nouvelle série de trois (figure 17.5).

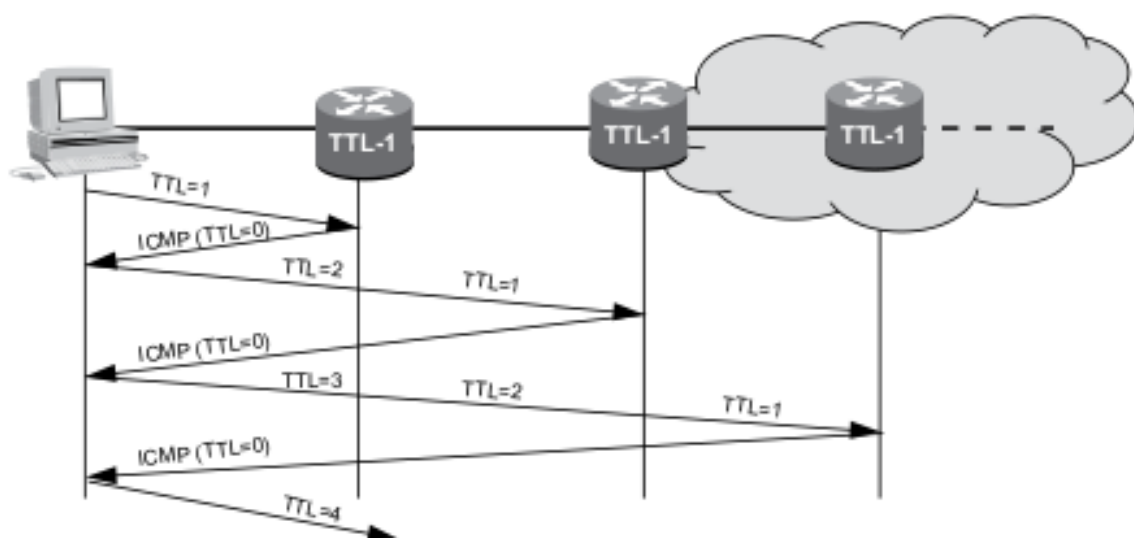


Figure 17.5 Principe de la détermination de route.

La première série est émise avec un TTL d'un, le premier nœud atteint décrémente le TTL, celui-ci étant alors à zéro, il élimine le datagramme et envoie à la source un message ICMP type 11 « *time exceeded* » (expiration de délai). Le premier nœud est ainsi découvert, l'utilitaire poursuit jusqu'à trouver la destination. Les requêtes Traceroute sont émises avec un

numéro de port > 30 000, le nœud destination retourne un message ICMP de type 3, code 3 « port inaccessible ».

C:> Tracert 195.221.126.186

Détermination de l'itinéraire vers cnam-st-martin.rap.prd.fr [195.221.126.186]
avec un maximum de 30 sauts :

```

 1  237 ms  156 ms  140 ms  nssta107.francetelecom.net [193.252.253.123]
 2  156 ms  156 ms  156 ms  GE2-201.ncidf103.Puteaux.francetelecom.net
 3  156 ms  156 ms  156 ms  pos6-0.nraub103.Aubervilliers.francetelecom.net [193.252.159.42]
 4  140 ms  156 ms  140 ms  pos13-3.ntaub201.Aubervilliers.francetelecom.net [193.252.103.18]
 5  156 ms  156 ms  156 ms  193.251.126.54
 6  140 ms  156 ms  140 ms  P14-0.PASCR2.Pastourelle.opentransit.net [193.251.128.105]
 7  156 ms  156 ms  155 ms  193.51.185.2
 8  140 ms  156 ms  140 ms  nri-b-pos11-0.cssi.renater.fr [193.51.179.10]
 9  156 ms  156 ms  156 ms  jussieu-pos4-0.cssi.renater.fr [193.51.180.157]
10  140 ms  156 ms  140 ms  193.50.20.73
11  156 ms  156 ms  155 ms  cr-cnam.rap.prd.fr [195.221.125.205]
12  140 ms  156 ms  156 ms  cnam-st-martin.rap.prd.fr [195.221.126.186]

```

Itinéraire déterminé.

Figure 17.6 Exemple de commande Traceroute.

La réponse à une commande Traceroute, illustrée figure 17.6, indique les différents nœuds traversés, de la machine source à la machine cible. Chaque ligne indique le rang du nœud traversé, le temps aller-retour (RTT) et, s'il est disponible, le nom du nœud et son adresse IP, cette information permet de déterminer la position géographique du nœud et éventuellement l'opérateur.

17.2 La résolution d'adresses

Les applications ne connaissent que l'adresse logique IP, alors que les données doivent être acheminées dans le réseau physique. À cette fin, la couche IP doit établir une correspondance entre l'adresse IP dans le réseau logique et l'adresse physique. Ce mécanisme s'appelle **mécanisme de résolution d'adresses**.

17.2.1 La résolution d'adresses dans les réseaux locaux IPv4

Dans les réseaux locaux, le protocole **ARP** (*Address Resolution Protocol*) permet à tout nœud du réseau d'obtenir l'adresse MAC de la machine destination. La figure 17.7 illustre ce mécanisme. En 1, le nœud d'origine,

d'adresse IP notée @IPa et d'adresse physique notée @Pha doit transmettre des données à la station d'adresse IP @IPc dont il ignore l'adresse physique. Il émet une demande de résolution d'adresses (ARP) sur le réseau physique à destination de la machine cible.

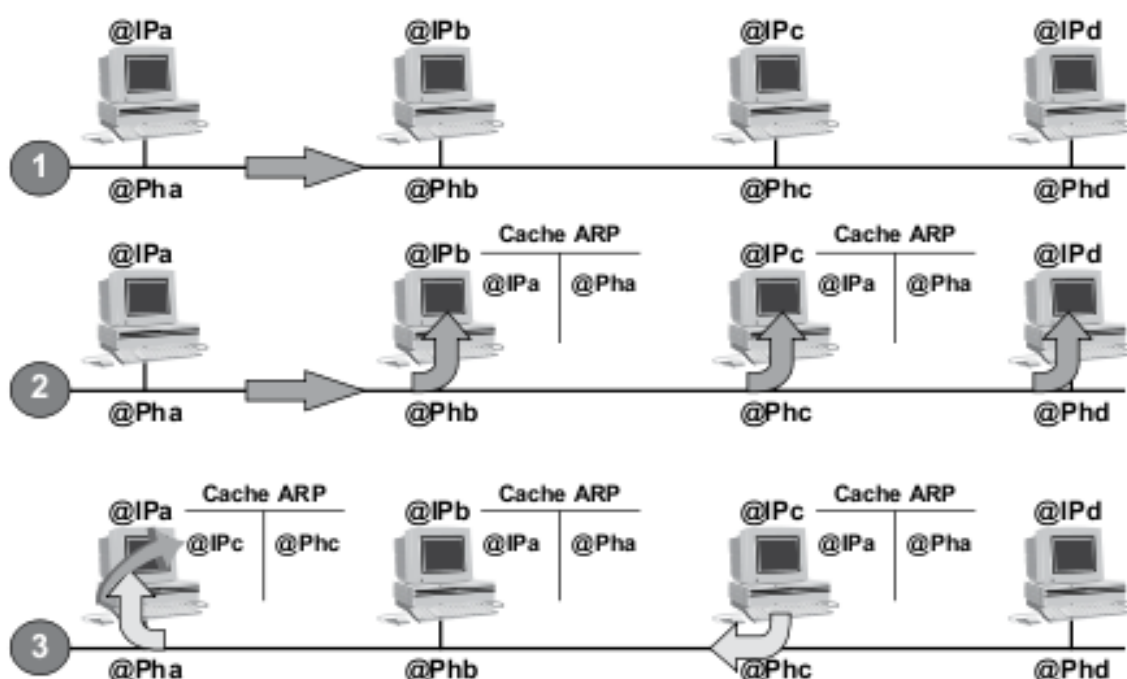


Figure 17.7 Le mécanisme de la résolution d'adresses.

Cette demande est encapsulée dans une trame MAC (Ethernet, Token Ring...) ; ne connaissant pas l'adresse du destinataire, le champ adresse destination est positionné à FF-FF-FF-FF-FF-FF (adresse de diffusion). En 2, tous les nœuds actifs sur le brin du réseau local, reconnaissant une demande ARP, en extraient l'adresse MAC origine et l'adresse IP origine et stockent ces valeurs, pour une éventuelle utilisation ultérieure, dans une table dite **cache ARP**. En 3, seule, la machine IPc qui a reconnu son adresse logique (@IPc) dans la requête répond en fournissant son adresse physique. Cette réponse est émise avec une adresse *unicast* puisque IPc connaît maintenant l'adresse physique de IPa. La réponse est stockée dans le cache ARP de la machine origine.

Pour limiter le trafic de résolution d'adresses et éventuellement détecter une duplication d'adresses, toute machine s'envoie, lors de sa mise sous tension, une demande ARP (*gratuitous ARP*). La figure 17.8 illustre le format d'un paquet ARP.

0	7	8	15	16	31
Espace d'adressage physique (Ethernet 0x0001)			Espace d'adressage logique (IP 0x0800)		
Longueur @Physique (Ethernet 6)		Longueur @Logique (IP 4)		Code opération (Requête 1, Réponse 2)	
Adresse physique émetteur					
@Phy émetteur (suite)			Adresse logique émetteur		
@Log émetteur (suite)			Adresse physique cible		
@Phy cible (suite)					
Adresse logique cible					

Figure 17.8 La structure du datagramme ARP.

À l'inverse du protocole ARP, le protocole **RARP** (*Reverse Address Resolution Protocol*) permet à une station qui ne dispose pas d'adresse IP (station sans disque, imprimante...) de s'en voir attribuer une. Le format du paquet RARP est identique à celui du protocole ARP. Seule, la valeur du champ *Code opération* est différente : 3 pour une requête et 4 pour la réponse. Une machine du réseau doit être configurée pour répondre à ces requêtes (serveur RARP). Le protocole RARP est aujourd'hui considéré comme obsolète.

17.2.2 La résolution d'adresses dans les réseaux locaux IPv6

À l'instar d'IPv4 une machine pour communiquer avec une autre sur un réseau local doit connaître l'adresse physique de cette dernière (adresse MAC). Cette adresse ne peut être déduite de l'adresse IP. En effet, l'adresse d'interface peut certes être construite à partir de l'adresse MAC mais elle peut aussi avoir été générée, pour des raisons de sécurité, aléatoirement par l'hôte distant. Il faut donc un protocole similaire à ARP pour découvrir cette adresse. La résolution d'adresses physiques dans IPv6 est réalisée par un mécanisme dit de découverte des voisins (RFC 2461, *Neighbor Discovery* ou ND) du protocole ICMPv6, messages de type = 135 (*Neighbor Solicitation*, NS) et 136 (*Neighbor Advertisement*, NA) représentés figure 17.9. Le mécanisme de découverte des voisins (ND) comprend un ensemble de fonctions qui déterminent les relations entre voisins, il remplace le protocole ARP, *Router Discovery* d'ICMPv4 et la Redirection d'ICMPv4). Le cache ARP est remplacé par un cache de voisinage.

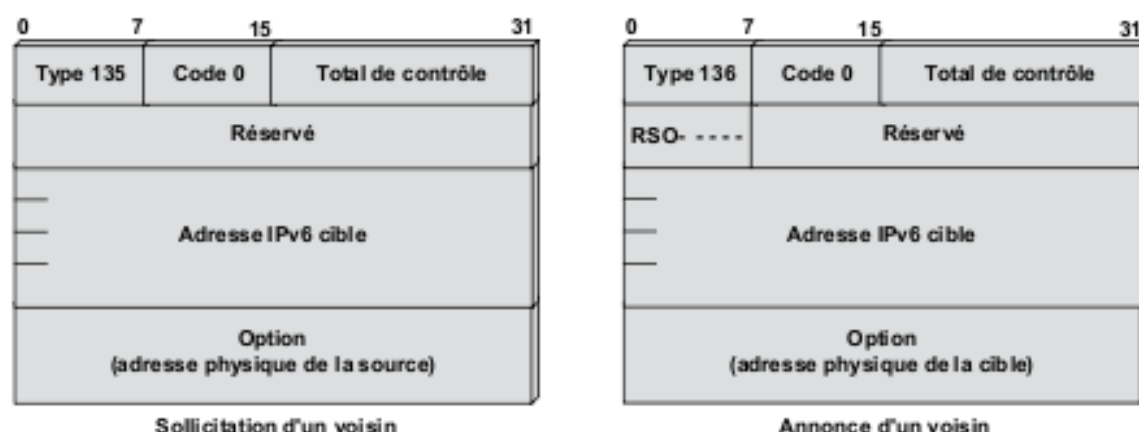


Figure 17.9 Messages de découvertes de voisins.

Le message de sollicitation d'un voisin permet d'acquérir les données d'un hôte situé sur le même lien physique ou à travers un système d'interconnexion de niveau 2 (commutateur ou pont). Le message est émis à l'adresse *multicast* sollicitée MAC-48 (33 :33:FF :00:00:01, tous les nœuds du lien local) ou l'adresse de groupe *multicast* du destinataire (33 :33:FF :03:04:05 dans l'exemple de la figure 17.10). De même au niveau IP, l'adresse IP peut être l'adresse de diffusion ou l'adresse IPv6 de la station sollicitée. Le champ **Adresse source** peut contenir l'adresse globale IPv6, l'adresse de lien local (voir l'exemple) ou encore une adresse non spécifiée. L'adresse **IPv6 cible** contient l'adresse IP de la machine questionnée.

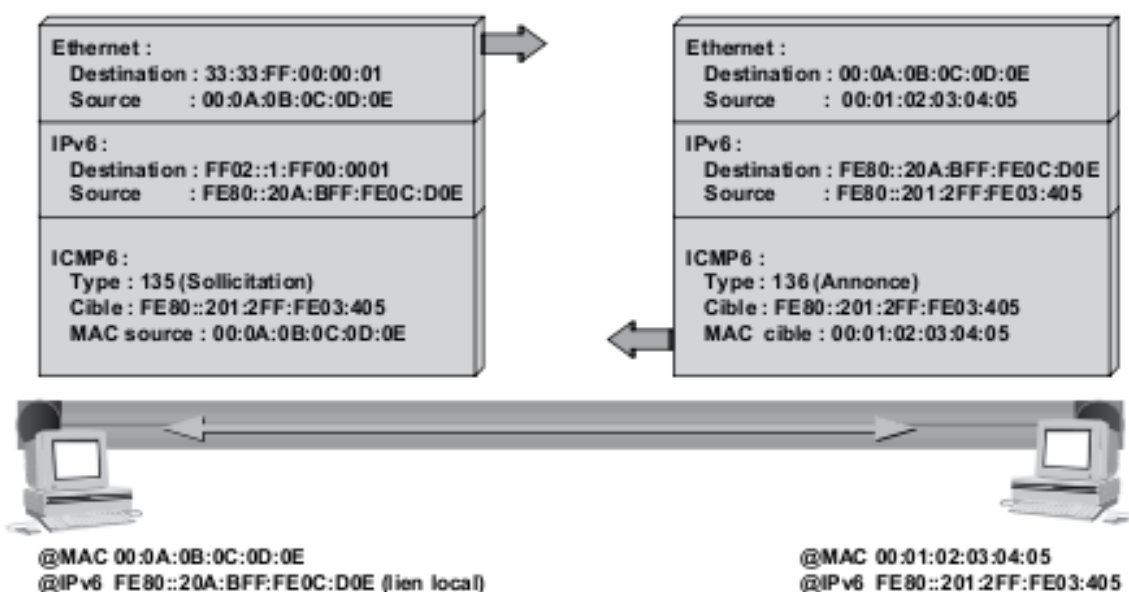


Figure 17.10 Résolution d'adresses physiques sous IPv6.

17.3 Les utilitaires de configuration dans IPv4

17.3.1 Généralités

Le protocole RARP permet à une station de se voir attribuer une adresse IP, les requêtes RARP utilisent une adresse de *broadcast* limitée. De ce fait, elles ne sont pas retransmises par les routeurs, ce qui implique un serveur RARP par réseau (sous-réseau). Le masque de sous-réseau peut-être découvert par une requête ICMP. Mais ce sont les seules informations que la machine peut obtenir. Le protocole **BOOTP** (*BOOTstrap Protocol*) permet à une machine sans disque de connaître l'intégralité des paramètres de configuration du réseau (adresse, masque, passerelle par défaut...). Le serveur BOOTP associe à l'adresse physique du client un profil. De ce fait, BOOTP n'est utilisable que dans un environnement relativement figé, ce qui en limite l'utilisation, notamment avec l'introduction du concept de mobilité (ordinateurs portables). **DHCP** (*Dynamic Host Configuration Protocol*), considéré comme une évolution de BOOTP, autorise une configuration automatique des stations et permet la mobilité.

17.3.2 DHCP

Reprenant les principes de BOOTP, DHCP (RFC 1541) offre de véritables fonctions de configuration automatique des stations. Contrairement à BOOTP qui affecte de manière statique une adresse IP, le serveur DHCP détient un jeu d'adresses valides et les paramètres associés de configuration IP à allouer dynamiquement aux clients DHCP. Les stations configurées par DHCP libèrent les adresses lorsqu'elles n'en ont plus besoin. DHCP propose trois mécanismes d'adressage :

- ▶ L'allocation manuelle, à l'instar de BOOTP ou RARP, DHCP alloue une adresse spécifique à un client. L'administrateur gère donc les adresses IP (DHCP statique) ;
- ▶ L'allocation automatique, elle permet, lors d'une configuration initiale, d'attribuer automatiquement à une station une adresse IP choisie par

le système parmi un pool d'adresses. La station conserve cette adresse tant qu'elle n'a pas été libérée explicitement par l'administrateur ;

- L'allocation dynamique, dans ce mode de fonctionnement, DHCP alloue temporairement (bail) une adresse IP. En fin de bail, la machine peut en demander le renouvellement.

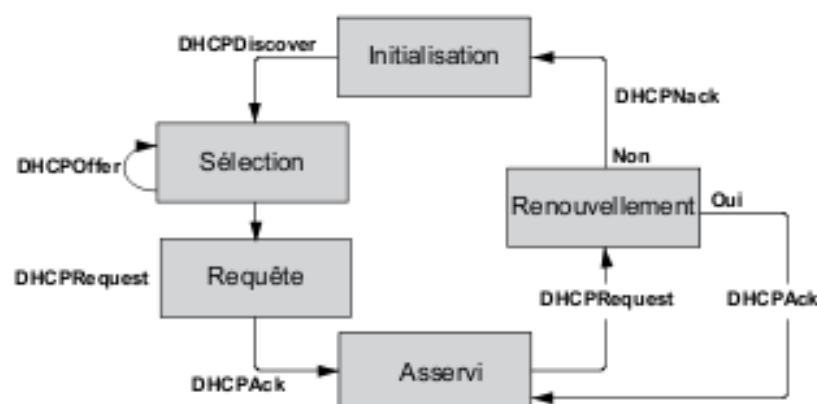


Figure 17.11 Les différents états d'un client DHCP.

La figure 17.11 décrit les différentes étapes d'une configuration automatique DHCP. La station, lors de son initialisation, diffuse un message d'exploration (**DHCPDiscover**). Tous les serveurs DHCP actifs sur le réseau formulent une offre de service (**DHCPOffer**). La station cliente passe alors dans l'état sélection ; lorsqu'elle a choisi un serveur DHCP (serveur élu), elle formule auprès de celui-ci une requête d'affectation d'adresse (**DHCPRequest**).

Le serveur DHCP sélectionné accuse réception de la requête (**DHCPAck**) en fournissant les éléments de configuration TCP/IP nécessaires à la station (adresse IP, masque de sous-réseau, adresse de la passerelle...) et la durée de validité de ces paramètres (bail). La station est alors dans l'état dit « asservi ». Si elle en a la possibilité, elle mémorise les informations. Elle pourra les utiliser lors des connexions futures jusqu'à expiration du bail, échéance au bout de laquelle elle devra formuler une demande de renouvellement (**DHCPRequest**). La requête initiale ou le renouvellement peut être refusé par le serveur élu (**DHCPNack**). La station est alors revenue à l'état initialisation.

Le message **DHCPRelease** permet au client de résilier son bail avant l'échéance de celui-ci. Le message **DHCPDecline** est utilisé par le client pour informer un serveur que son offre est invalide. Enfin, le message **DHCPInform** permet à une machine d'obtenir des paramètres de configuration supplémentaires.

DHCP utilise les mêmes ports (67 pour le serveur, 68 pour le client) et le même format de message que BOOTP, seul le dernier champ diffère.

17.4 L'auto-configuration dans IPv6

17.4.1 Généralités

L'auto-configuration est un mécanisme qui permet à tout équipement IPv6 d'acquérir automatiquement les paramètres nécessaires à son fonctionnement (adresse IP, routeur par défaut, DNS...). IPv6 offre deux mécanismes d'auto-configuration :

- ▶ auto-configuration avec état (*Stateful*) reposant sur la version IPv6 du protocole DHCP (DHCPv6, RFC 3315) ;
- ▶ auto-configuration sans état (*Stateless Address Autoconfiguration*), s'applique à tous les équipements terminaux, sauf les routeurs.

17.4.2 L'autoconfiguration sans état

L'autoconfiguration sans état repose sur le mécanisme dit de découverte des voisins (RFC 2461, *Neighbor Discovery* ou ND) du protocole ICMPv6, messages de type = 133 (*Router Solicitation*, RS) et 134 (*Router Advertisement*, RA) représentés figure 17.12.

Le message RS est similaire au message de sollicitation de voisins. Il est émis à l'intention du groupe *multicast* « FF02 ::2 » (tous les routeurs). La source qui ne dispose pas encore d'adresse IP peut utiliser l'adresse non spécifiée (« :: »). Le message de réponse (RA) est émis périodiquement, ou en réponse à un message RS à l'adresse du groupe *multicast* « FF02 ::1 » (tous les nœuds du lien local) ou à l'adresse de la station qui a émis le message RS.

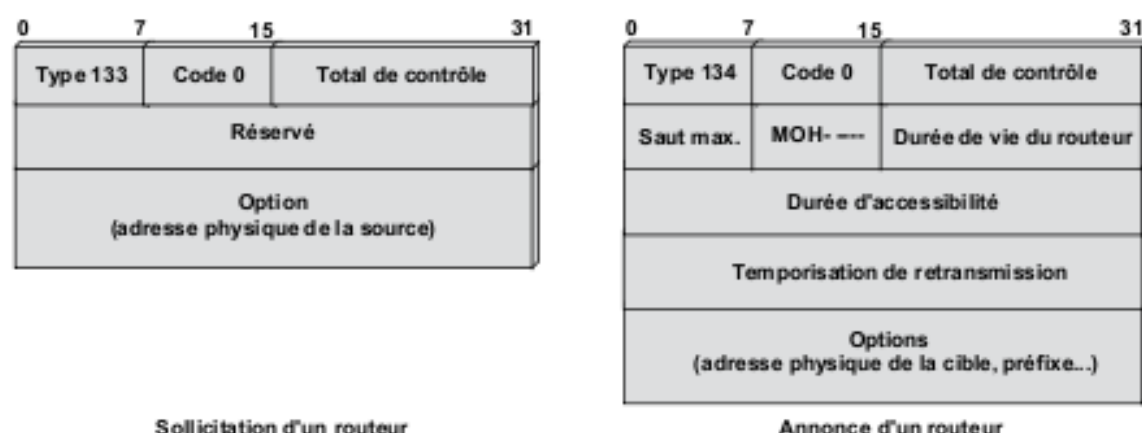


Figure 17.12 Messages RS et RA.

Le champ **Saut max**, non nul, spécifie la valeur d'initialisation du champ nombre de sauts des datagrammes IPv6. Le champ **Durée du routeur** indique, en seconde, la durée pendant laquelle le routeur répondant fait fonction de routeur par défaut, à « 0 » cette valeur indique que ce routeur ne remplit pas les fonctions de routeur par défaut. Le champ **Durée d'accessibilité** spécifie la durée de conservation (validité) des informations du cache. Le champ **Temporisation de retransmission** indique la périodicité d'émission des messages RA (de 4 à 1 800 secondes). Le **champ options** comporte le ou les préfixes utilisés sur le lien local, l'adresse physique du routeur annonçant (MAC-48) et, la MTU. Un champ de bits complète ces informations.

17.5 IP et la mobilité

17.5.1 Généralités

La notion de mobilité est apparue avec l'avènement des ordinateurs portables. Cependant, si ces ordinateurs permettent de se déplacer, il restait à résoudre le problème d'accès aux applications auxquelles il était nécessaire de masquer le déplacement du mobile, c'est ce problème que résout l'**IP Mobile** (RFC 2002, *IP Mobility Support*).

La mobilité recouvre deux notions distinctes. La première consiste à accéder à ses données depuis un point d'accès fixe d'un réseau étranger au

réseau d'appartenance. L'ordinateur a été déplacé mais reste fixe durant la connexion, on parle alors d'**ordinateurs nomades** ou de macromobilité. La seconde, rendue possible par l'apparition des réseaux hertziens autorise l'ordinateur à rester en contact avec son réseau d'appartenance alors qu'il se déplace et éventuellement change de réseau d'accueil, on parle alors d'**ordinateurs mobiles**. Bien que la dénomination du protocole soit IP mobile, *IP Mobility Support* est mal adapté aux changements rapides de réseau d'accueil.

L'IP mobile est un enrichissement du protocole IP pour permettre à un ordinateur mobile (nomade) d'accéder à ses applications depuis n'importe quel réseau d'accueil sans qu'il soit besoin de le reconfigurer : l'ordinateur mobile communique en utilisant son adresse IP d'origine.

17.5.2 Le principe de la mobilité sous IPv4

L'agent mobile possède une adresse principale dite encore adresse mère (*Home address*, @IPh) dans son réseau d'origine (réseau de domiciliation ou *Home network*) et se voit attribuer une adresse temporaire (*Care_of address*, @IPv) dans le réseau visité (*Foreign network*). La mobilité devant être transparente aux applications, le mobile utilise son adresse principale pour communiquer, par conséquent toutes les réponses sont adressées à son réseau d'origine. Dans ce réseau (*Home network*), un routeur IP, doté d'un agent IP mobile (agent domestique, ou *Home agent*), intercepte les messages à destination du mobile et lui réexpédie à son adresse temporaire (*Care_of address*). On distingue deux modes de réexpédition. Le premier dit réexpédition par agent extérieur nécessite dans le réseau visité la présence d'un agent relais (*Foreign agent*), c'est cet agent relais qui attribue l'adresse temporaire au mobile. La figure 17.13 illustre ce mode de fonctionnement.

L'échange dans le sens mobile vers son correspondant distant ne pose aucun problème particulier ; l'échange utilise l'adresse d'origine du mobile (source : @IPh) et l'adresse du correspondant distant (destination : @IPc). La mobilité étant transparente à l'ordinateur distant, celui-ci adresse sa réponse à l'adresse IP du mobile dans son réseau de domiciliation (destination : @IPh). Ces messages sont interceptés par l'agent domestique et

réexpédiés vers le réseau visité. Pour cela, le message IP reçu du correspondant distant est encapsulé dans un datagramme IP (IP dans IP, RFC 2003) dont les adresses correspondent à celles des deux agents (domestique et relais). L'ensemble forme un tunnel IP.

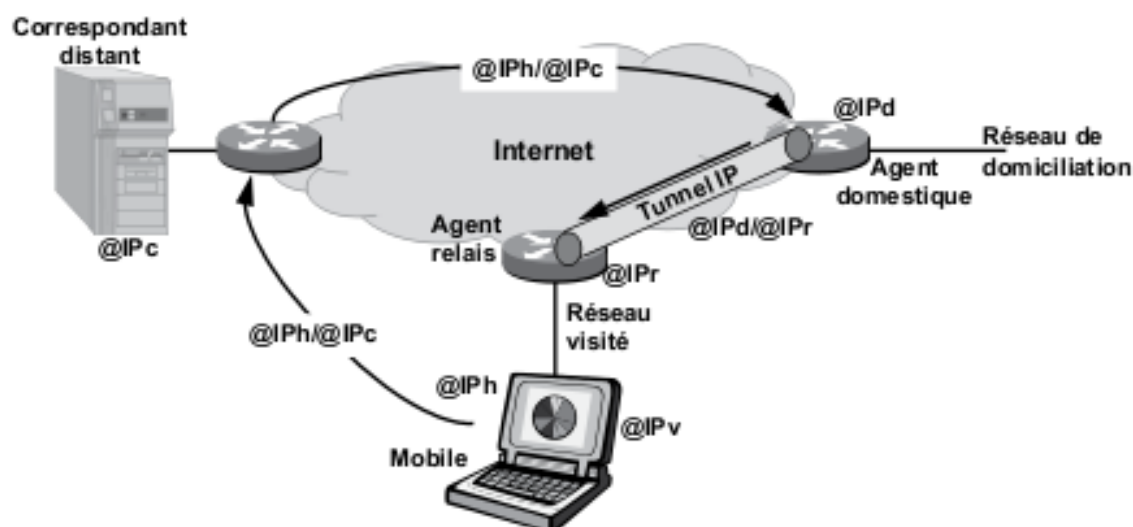


Figure 17.13 La mobilité d'IPv4 (10.31).

17.5.3 La mobilité dans IPv6

La mobilité IPv6 est une évolution de la mobilité IPv4 qui, profitant des nouvelles fonctionnalités d'IP, simplifie et optimise la mobilité. L'agent relais (*Foreign agent*) disparaît, seul le mode réexpédition par colocataire subsiste.

Dans la mobilité IPv6, un mécanisme spécifique permet à chaque nœud en communication avec le mobile d'apprendre et de mémoriser l'adresse temporaire du mobile (association adresse mère/adresse temporaire). Lorsqu'un nœud doit adresser un datagramme au mobile, s'il possède une association pour l'adresse mère du mobile, il utilise l'option IPv6 « en-tête de routage » pour adresser directement le datagramme au mobile dans le réseau visité (optimisation du routage), sinon il envoie le datagramme vers le réseau mère, ce datagramme intercepté par l'agent domestique est réacheminé par encapsulation vers le réseau visité (figure 17.14).

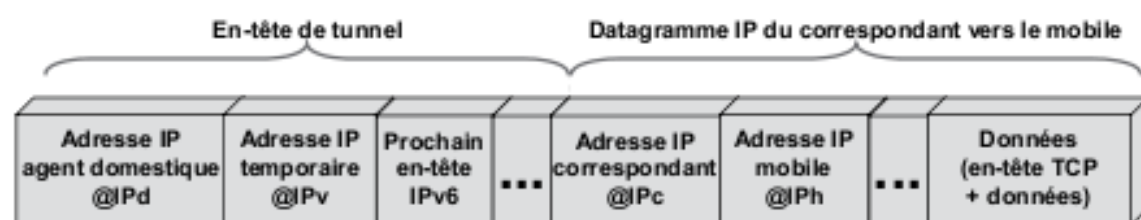


Figure 17.14 L'encapsulation IPv6 dans IPv6.

18

Les applications de l'environnement TCP

Contrairement à ce qui est fréquemment écrit, TCP/IP ne possède pas de couche application, mais des applications indépendantes les unes des autres. Celles-ci s'appuient directement sur TCP ou UDP et non sur des services de niveau application comme le font les applications du modèle de référence. Les paragraphes qui suivent n'ont pas pour objectif de décrire toutes les applications définies par un RFC mais seulement quelques-unes d'entre elles parmi les plus utilisées.

18.1 Notions d'annuaire

Un annuaire est un système permettant, à partir d'une information, d'en connaître une autre. L'annuaire téléphonique fournit, à partir d'un nom, le numéro de téléphone et généralement l'adresse de la personne concernée. Un annuaire électronique est une base de données consultable par une application, il délivre à cette dernière les informations nécessaires à son bon fonctionnement. Par exemple, une application peut consulter un annuaire pour y connaître les droits d'un utilisateur.

18.1.1 Le service de noms (DNS)

■ Principe

Rappelons que le nommage est une notion complémentaire de celle de l'adressage, l'un désigne l'objet (objet nommé) l'autre sa localisation. Dans l'environnement TCP/IP, le service de noms consiste simplement à associer un nom à une adresse IP (résolution de noms). Ce service peut être rendu simplement par lecture d'un fichier local (etc/host) ou par consul-

tation d'un annuaire électronique : le **DNS**¹, base de données distribuées s'appuyant sur les protocoles de transport UDP et TCP (messages de plus de 512 octets). Le port 53 est utilisé dans les deux cas (TCP et UDP).

D'origine **IAB** (*Internet Activities Board*), le DNS repose sur le modèle relationnel client/serveur. La partie cliente, le solveur (*resolver*), est chargée de résoudre la correspondance entre le nom symbolique de l'objet et son adresse réseau.

■ L'espace de nommage

Les noms sont organisés selon une structure arborescente hiérarchique (arbre inversé) appelée **espace de nommage** (figure 18.1). La racine est au sommet, son nom de domaine est vide, elle est symbolisée par un point « • ». Le nombre de niveaux est limité à 127. Un nom ne peut dépasser 255 caractères et chaque niveau est limité à 63 caractères. S'ils sont dans des domaines différents, des nœuds ou des feuilles peuvent avoir des noms identiques.

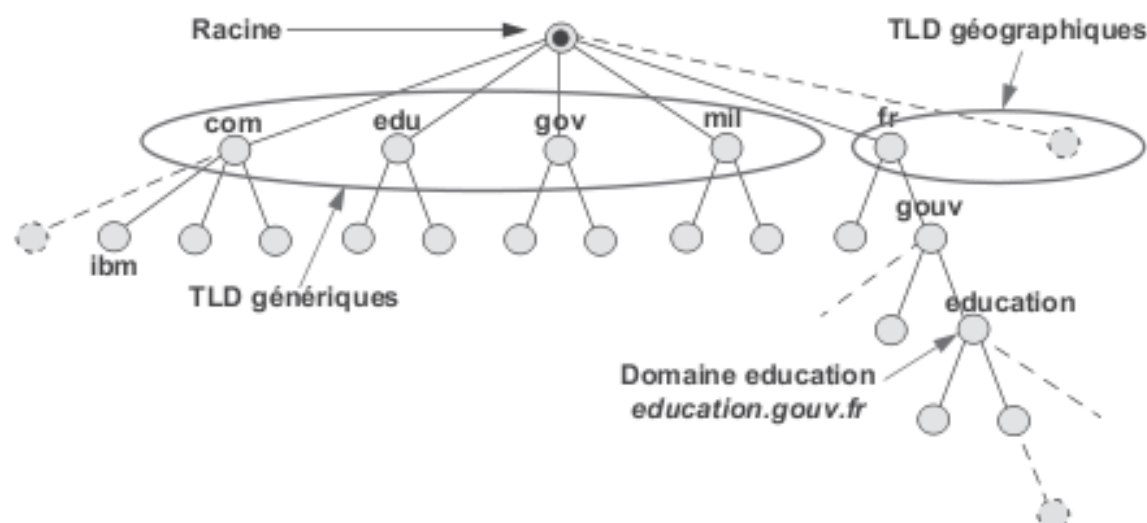


Figure 18.1 La structure de l'espace de nommage.

1 De façon plus stricte : DNS (Domaine Name System) désigne l'ensemble des organismes qui gèrent les noms de domaine, DNS (Domain Name Service) correspond au protocole d'échange d'information entre les domaines, enfin DNS (Domain Name Server) représente le serveur de noms.

Les domaines de premier niveau de la hiérarchie de la figure 18.1 ou **TLD** (*Top Level Domain*), gérés par l'**ICANN** (*Internet Corporation for Assigned Names and Numbers*) sont regroupés en deux ensembles : les **gTLD** (*generic TLD*, TLD génériques) et les **ccTLD** (*country code TLD*, TLD géographique). Les ccTLD au nombre de 248 sont codés sur deux lettres suivant le code ISO 3166. La gestion en est confiée à des organismes régionaux (*Registry*) et la commercialisation à des organismes privés (*Registrar*). En France c'est AFNIC (Association française pour le nommage Internet en coopération) qui gère l'espace de nom « .fr ». Aucune signification n'est imposée aux noms de domaine, sauf pour le premier niveau (tableau 18.1).

Tableau 18.1 Les domaines génériques (gTLD).

gTLD	Signification	Destination / ouverture	Date d'ouverture
.aero	aéronautique	entreprises du secteur aéronautique	2002
.biz	business	entreprises commerciales (à tous)	2001
.cata	catalan	communauté linguistique et culturelle catalane	2006
.com	commercial	organisations commerciales possédant, en principe, des implantations sur plusieurs domaines géographiques (à tous)	1995
.coop	coopérative	pour les coopératives	2002
.edu	éducation	établissements d'enseignement supérieurs et universités	1995
.gov	gouvernement	établissements gouvernementaux des États-Unis	1995
.info	information(s)	à tous	2001
.int	international	organisations internationales reconnues par des traités internationaux	1998
.jobs	jobs	responsables de ressources humaines	2005
.mobi	mobiles	à tous	2006
.name	nom de famille	aux personnes physiques et aux personnes morales pour la protection de leur marque	2002
.net	réseau	à tous	1995
.mil	militaire	organisations militaires américaines	1995

gTLD	Signification	Destination / ouverture	Date d'ouverture
.museum	musée	pour les musées répondant à la définition de l'International Council of Museum (ICOM)	2001
.org	organisation/ association	organisations non commerciales et non gouvernementales à but non lucratif (à tous)	1995
.pro	professionnel libéral	pour les professions libérales (avocats, médecins...)	2002
.travel	voyage/ tourisme	aux professionnels du tourisme	2005
.arpa		domaine réservé à la résolution de nom inverse	1995

Dans le réseau Internet, la racine (*root*) est l'élément de base de toute la hiérarchie sur laquelle repose le fonctionnement d'Internet. Aussi la gestion des TLD de premier niveau est-elle répartie sur 13 serveurs répartis à travers le monde. Ces serveurs appartiennent au domaine **root-servers.net**. Le serveur maître A.root-server.net (États-Unis) est géré par VeriSign Global Registry Services, les autres serveurs sont des serveurs miroirs désignés sous les noms de B.root-servers.net, C.root-servers.net... K.root-servers.net.

■ L'anycast et le DNS

L'*anycast* est une technique qui autorise plusieurs machines géographiquement dispersées à utiliser une même adresse IP. Toute requête à destination d'une adresse *anycast* est acheminée vers le système le plus proche au sens routage du terme. Cette technique utilisée pour les serveurs DNS (RFC 3258) autorise la duplication de ces serveurs. Ainsi, pour un même serveur DNS racine, plusieurs serveurs physiques peuvent être mutualisés. Cette répartition diminue la charge de chaque machine physique, accroît les performances et augmente la résistance aux attaques du type déni de service¹. La technique d'*anycast* rend obso-

¹ Une attaque en déni de service (DoS) consiste à surcharger un système pour que celui-ci s'effondre. C'est suite à une attaque de ce type (2002) que la technique d'*anycast* a été mise en œuvre dans les serveurs DNS d'Internet.

lète la notion de localisation géographique. Les 13 serveurs racines de l'Internet correspondent en fait à 120 serveurs physiques dont deux sont hébergés par la société SFINX en France.

■ La résolution de nom

□ Principe

Un nœud est désigné par l'arborescence complète de son nom (**FQDN**, *Fully Qualified Domain Name*). La résolution est l'opération qui consiste à mettre en relation le nom du nœud et son adresse IP. Le client DNS ou solveur (*Resolver*) est un programme de type *daemon*. Sur sollicitation d'un programme demandeur, il est chargé d'émettre les demandes et de traduire les réponses. Lors de la configuration d'une station IP, on lui fournit son nom de domaine, l'adresse de son serveur local de noms et, éventuellement, une liste ordonnée de serveurs de noms.

Le client solveur (figure 18.2) interroge le serveur de noms local. Si la recherche est infructueuse le serveur local interroge le serveur de niveau supérieur (recherche récursive) ou le client interroge lui-même d'autres serveurs (requête itérative).

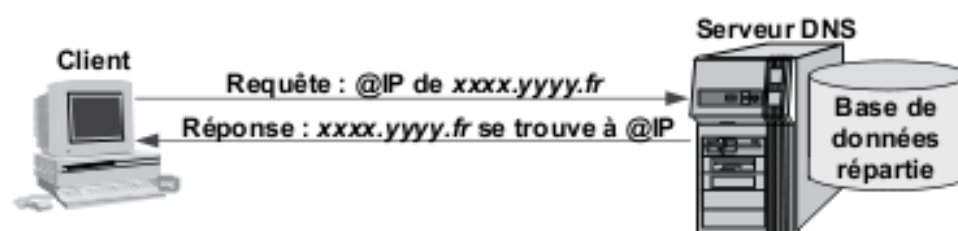


Figure 18.2 La résolution de nom.

En principe, un client émet des requêtes de type récursif, charge au serveur d'émettre, si besoin, des requêtes de type itératif pour rechercher la réponse.

□ Résolution inverse

À l'instar de la résolution d'adresses (RARP), la résolution de noms inverse permet d'obtenir, à partir de l'adresse IP, le nom de la machine. A cet effet, le domaine *arpa*, sous-domaines *in-addr* (IPv4) et *ip6* (IPv6), a été prévu (figure 18.3).

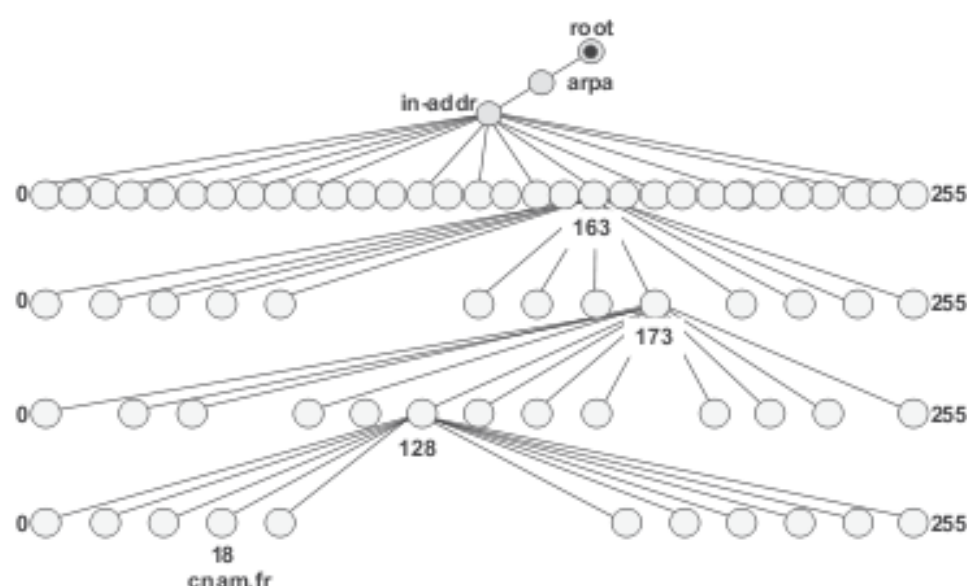


Figure 18.3 L'arbre de résolution inverse.

L'arbre inverse considère l'adresse comme un nom : par exemple l'adresse du Conservatoire national des Arts et Métiers (163.173.128.18) est traduite par :

« 18.128.173.163.in-addr.arpa »

À chaque octet de l'adresse IP correspond un nœud de l'arbre. Chaque sous-domaine ainsi défini comporte 256 sous-domaines. Le quatrième niveau correspond au nom du serveur connaissant le nom de domaine associé à cette adresse. La notation retenue pour IPv6 est :

« 1.0.0.0.1.0.0.0.0.0.0.0.0.0.0.0.1.0.0.0.6.0.0.3.0.6.6.0.1.0.0.2.ip6.arpa »

pour l'adresse « 2001:660 :3006:1::1:1:l » (adresse de ns3.nic.fr).

□ UDP ou TCP ?

Le protocole DNS utilise indifféremment le transport TCP ou UDP et, dans les deux cas, le port 53. Les messages UDP sont limités à 512 octets (hors en-tête). Sous UDP, lorsque la réponse dépasse 512 octets, celle-ci est envoyée au client tronquée à 512 octets avec le bit **TC** positionné à 1 (*TronCation*, message tronqué). En principe, le résolveur émet alors une nouvelle requête sur TCP. Il est possible de forcer l'utilisation de TCP,

mais ceci au détriment des performances. Cependant, TCP est toujours utilisé pour effectuer les transferts de zone.

18.2 Le transfert de fichiers

Le transfert de fichiers est l'une des applications les plus utilisées sur les réseaux. Le modèle TCP/IP en décline deux versions. L'une allégée (**TFTP**, *Trivial File Transfer Protocol*) nécessite peu de mémoire et peut tenir en mémoire morte des machines sans disque (terminal X, par exemple) et ainsi permettre le téléchargement du système. TFTP s'appuie sur UDP. L'autre version, **FTP** (*File Transfer Protocol*) constitue un véritable système de manipulation de fichiers à distance.

18.2.1 TFTP (Trivial File Transfer Protocol, RFC 1350)

TFTP permet le transfert de données en lecture (**RRQ**, *Read ReQuest*) ou en écriture (**WRQ**, *Write ReQuest*) de fichiers en ASCII (mode dit *netascii*) ou en flux d'octets (mode dit *octet*). Le transfert a lieu par bloc de 512 octets numérotés. La fin du transfert est indiquée par un message de données dont la longueur est inférieure à 512 octets (figure 18.4).

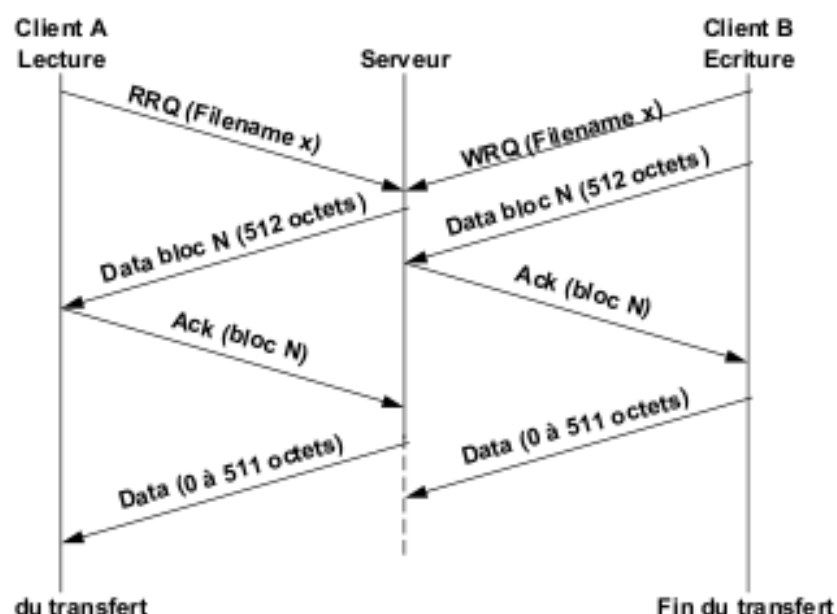


Figure 18.4 L'échange TFTP.

TFTP s'appuie sur UDP, c'est donc à TFTP de gérer les paquets perdus. Le protocole est du type *Send and Wait* (émettre et attendre), chaque extrémité gérant une reprise sur temporisation (transmission symétrique). En principe, un message d'erreur peut se substituer à un ACK, une retransmission est alors réalisée. Cependant, dans la plupart des implémentations, un message d'erreur provoque l'arrêt du transfert.

18.2.2 FTP (File Transfert Protocol, RFC 959)

L'originalité de FTP est d'ouvrir pour chaque session FTP deux connexions simultanées. L'une sur le port 21 (FTP), l'autre sur le port 20 (FTP_Data). La première connexion, connexion de contrôle ou de service, sert à l'échange des messages FTP (connexion de signalisation), l'autre au transfert de données (figure 18.5). La demande de connexion FTP est établie sur le port 21 et reste active durant toute la session FTP (connexion permanente). La connexion de transfert sur le port 20 n'est active que durant le transfert effectif d'un fichier (connexion temporaire).

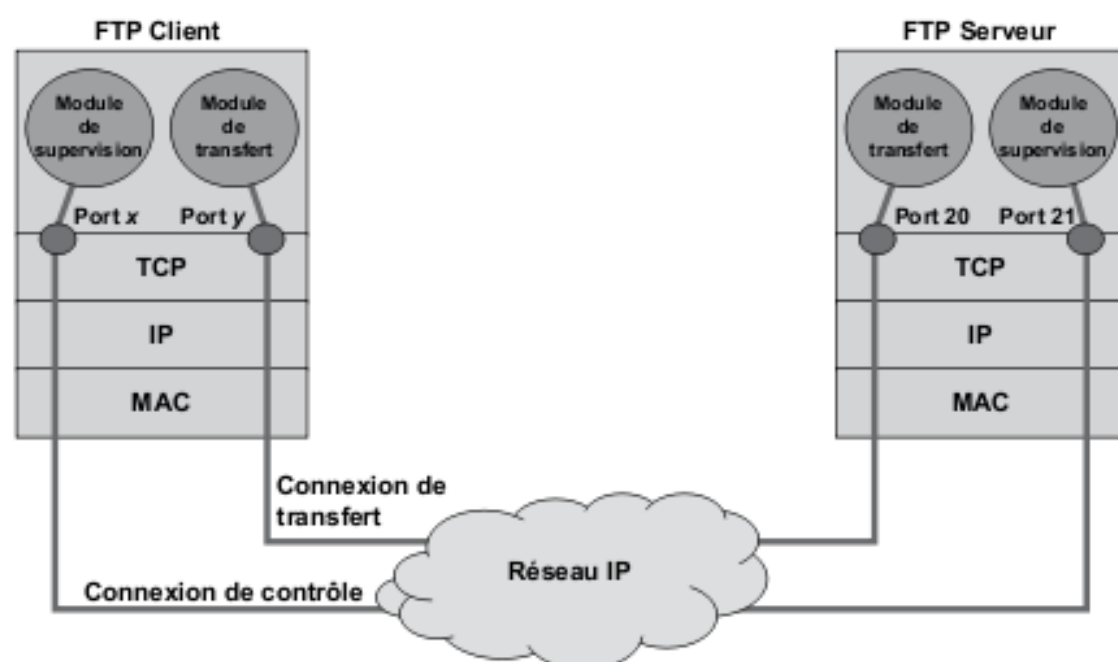


Figure 18.5 Les deux connexions de FTP.

Contrairement à TFTP, FTP réalise un contrôle d'accès avant l'acceptation de toute connexion. En principe, il faut posséder un compte sur le serveur

FTP pour pouvoir s'y connecter. La session FTP commence par une procédure de type « *login* » : nom d'utilisateur, mot de passe. Il est aussi possible de se connecter sans compte (invité). Dans ce cas, le nom d'utilisateur est *anonymous* et le mot de passe *guest* (connexion FTP anonyme).

18.3 L'émulation de terminal (Telnet)

Le terminal virtuel est un logiciel en mode client/serveur. Le terminal serveur émule vis-à-vis d'un applicatif sur la machine serveur un terminal local avec lequel l'application échange des messages, tandis que sur une machine distante, cliente de l'application, un terminal dit client est émulé sur un terminal physique. L'échange se fait en deux temps. Le premier a lieu entre le terminal client réel et le terminal serveur virtuel, le second entre le terminal serveur virtuel et l'application. Cette approche est illustrée figure 18.6.

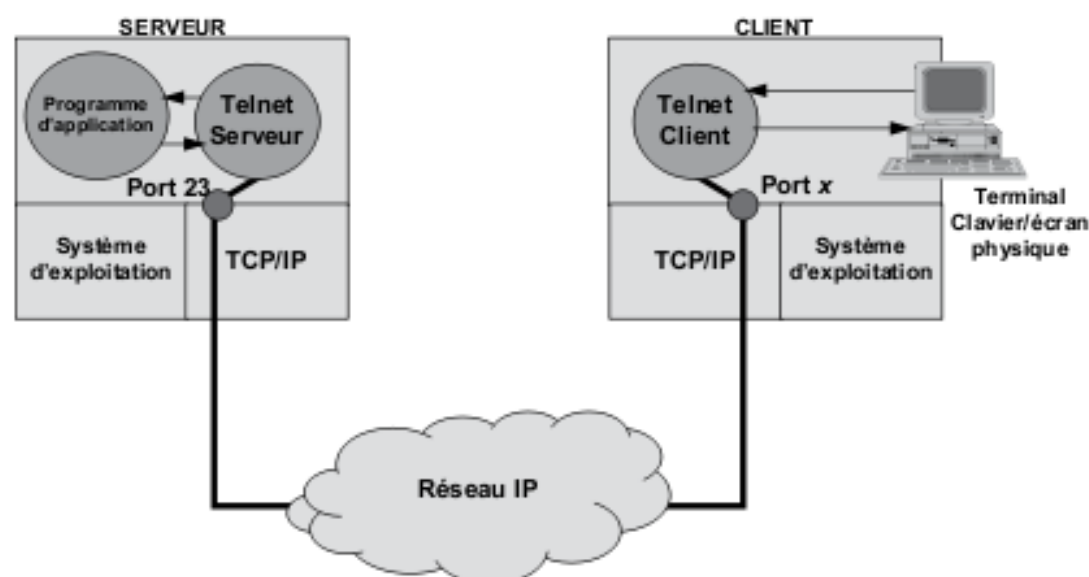


Figure 18.6 Principe du terminal virtuel.

Le terminal Telnet ne permet pas seulement de se connecter pour une session de travail sur une machine distante, il a longtemps été utilisé comme terminal de consultation de données sur différents serveurs (météo...). Le déploiement d'Internet et des navigateurs rend cette dernière utilisation obsolète.

18.4 La messagerie électronique

18.4.1 Introduction

Apparues dans les années 1970 sur les grands systèmes informatiques, les messageries électroniques autorisent une communication interpersonnelle sans mise en relation directe des correspondants (communication asynchrone). Les messages (**eMail**, *Electronic Mail* ou courriel en français pour courrier électronique) sont déposés dans une boîte aux lettres (**BAL**) départ et consultés dans une boîte aux lettres arrivée (figure 18.7).

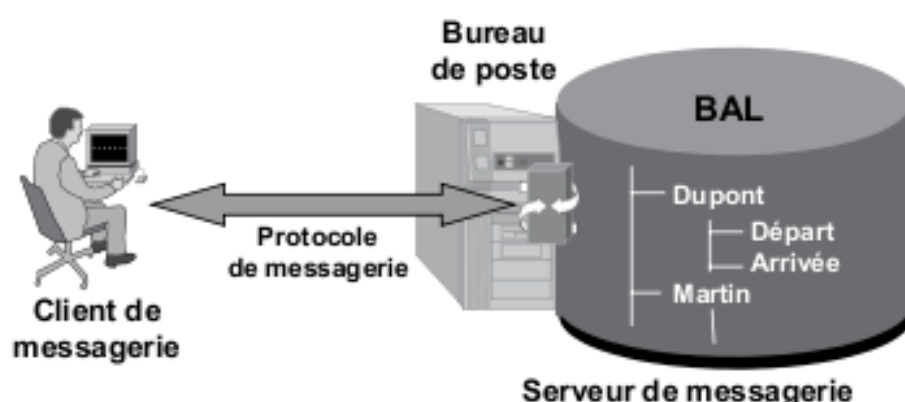


Figure 18.7 Principe d'un système de messagerie.

Chaque utilisateur dispose de ses propres boîtes aux lettres départ et arrivée, où sont mémorisés les messages reçus et envoyés. Indépendamment de la création, de l'envoi et de la réception des messages un service de messagerie électronique offre les services suivants :

- ▶ l'accusé de réception (l'expéditeur est averti du bon acheminement de son message) ;
- ▶ l'accusé de prise de connaissance (l'expéditeur est informé de l'ouverture du document) ;
- ▶ l'archivage des messages reçus ou envoyés ;
- ▶ la gestion d'un historique des messages ;
- ▶ l'envoi de pièces jointes (fichiers textes, documents multimédias, documents tableurs, fichiers binaires...). Cette dernière fonctionnalité

transforme un système de messagerie en système de transfert de fichiers ;

- ▶ la gestion d'une liste de correspondants (annuaire) ;
- ▶ la création de listes de diffusion.

Outil de communication, la messagerie électronique est devenue le support de toutes les applications nécessitant un besoin de collaboration (rapport rédigé en commun, outils de planification...) et de coordination (agenda...) donnant naissance à un ensemble d'applications orientées vers le travail en groupe (*groupware*).

18.4.2 Architecture du système de messagerie

Un système de messagerie¹ (**MHS**, *Message Handling System*) est essentiellement composé d'un réseau de transport des messages (**MTS**, *Message Transfer System*), des unités d'accès au système (**MUA**, *Message User Agent* ou simplement **UA**) et d'acheminement (**MTA**, *Message Transfer Agent*). L'interface avec laquelle l'utilisateur (**UI**, *User Interface*) prépare et envoie son message n'est pas incluse dans le système de messagerie. Lorsque l'UA ne dispose d'aucune possibilité de stockage, ou que celle-ci est inaccessible, un système de mémorisation intermédiaire peut être défini (**MS**, *Message Storage*).

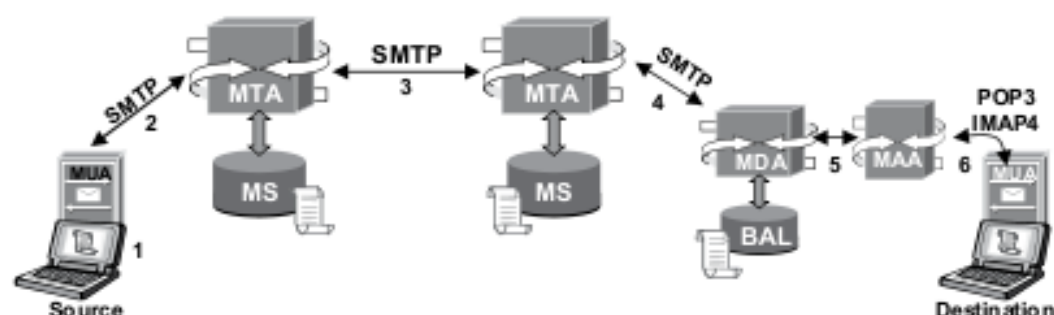


Figure 18.8 Architecture d'un système de messagerie.

Lorsqu'un utilisateur souhaite envoyer un message :

12. Il le rédige à l'aide de son client de messagerie (UI et MUA). Le protocole **MIME** (*Multipurpose Internet Mail Extensions*) permet le codage,

¹ La terminologie est celle définie par la messagerie UIT X.400

la mise en forme du message (attributs : gras, souligné...) et de joindre diverses pièces jointes selon différents formats (textes, photographies, vidéo...);

13. Le message est transmis au MTA de l'utilisateur (serveur de messagerie SMTP);
14. le message pris en charge par le MTS est envoyé au MTA du destinataire (protocole SMTP);
15. le serveur transmet le message à l'agent de distribution (**MDA**, *Message Delivery Agent*) qui le mémorise dans la BAL du destinataire;
16. l'agent d'accès (**MAA**, *Message Access Agent*) sur sollicitation d'un protocole de relève (POP3, *Post Office Protocol*, ou IMAP4, *Internet Message Protocol*) transfère le message sur le client de messagerie destinataire (mode *offline*). Il est aussi possible à l'aide d'un logiciel *Webmail* de gérer directement les messages sur le serveur de messagerie (mode *online*);
17. Enfin (mode *offline*), le destinataire consulte son message via son client de messagerie (MUA).

18

18.4.3 La diffusion des messages

La fonction de routage est assurée dans le MTS (*Message Transfer System*) après mise en relation de son adresse de messagerie (adresse globale) avec l'adresse IP du serveur de messagerie du domaine géré par le fournisseur d'accès. Une adresse de messagerie (RFC 5321 et 5322) respecte le format suivant :

Nom_Boîte_Aux_Lettres @ Nom_Domaine_DNS_Du_Serveur
 Nom_Boîte_Aux_Lettres @ @IP_Serveur
 claudeservin@provider.fr

Les adresses sont gérées localement par un carnet d'adresses qui autorise des adresses de groupes. Un message peut contenir plusieurs destinataires, une seule instance du message est transmise, c'est le serveur MTA d'accueil qui dupliquera le message en autant de messages que de destinataires. Ces destinataires peuvent être :

- ▶ des destinataires principaux ;
- ▶ des destinataires pour copie (*Cc, Carbon copy*) ;
- ▶ des destinataires en copie cachée invisibles des autres destinataires (*Cci, Carbon Copy invisible* ou *Bcc, Blind Carbon Copy*).

Les adresses multiples doivent être séparées par un point-virgule suivi d'un espace.

18.4.4 Les protocoles de messagerie

Le protocole **SMTP** (*Simple Mail Transfer Protocol*, RFC 821), assure l'acheminement des messages en réalisant une connexion en mode point à point entre serveurs MTA (mode connecté sur TCP, port 25). Le courrier est remis directement au serveur de courriers du destinataire.

Le protocole **POP3** (*Post Office Protocol*, RFC 1939) mis en œuvre par les clients de messagerie récupère les messages sur le serveur (serveur POP). Les messages récupérés sont effacés, mais présents sur la machine du destinataire où ils peuvent être consultés, hors connexion, jusqu'à leur effacement. De nombreux clients de messagerie n'effacent les messages transférés qu'après un délai déterminé (configuration du logiciel client).

Le protocole **IMAP4** (*Internet Message Access Protocol*, RFC 1733 et 2060) est un protocole alternatif au protocole POP3. IMAP ne transfère pas les messages, la gestion de ceux-ci se fait directement sur le serveur IMAP : lecture, effacement... Il autorise ainsi l'accès aux services de messagerie depuis n'importe quelle machine.

18.5 Les notions de middleware

18.5.1 Définitions

Le *middleware* peut être vu comme un complément de services réseau autorisant un dialogue entre applications de type client/serveur (figure 18.9). Sous cette appellation, on désigne généralement les différents outils de développement qui permettent d'offrir aux applications une interface d'accès aux différents services du réseau (**API**, *Application Programming*

Interface) ; il est le lien entre le *hardware* (l'environnement réseau) et le *software* (les applications).

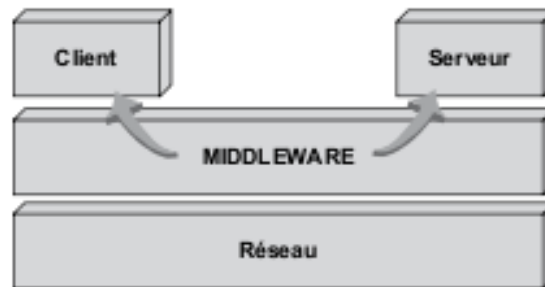


Figure 18.9 Position du middleware.

18.5.2 Les exemples d'outils middleware dans TCP/IP

■ RPC (Remote Procedure Call)

En programmation traditionnelle, les programmes et les procédures sont liés pour ne former qu'un seul programme. Le concept **RPC** (procédures éloignées) suppose que le programme principal et les procédures forment des programmes séparés (figure 18.10).

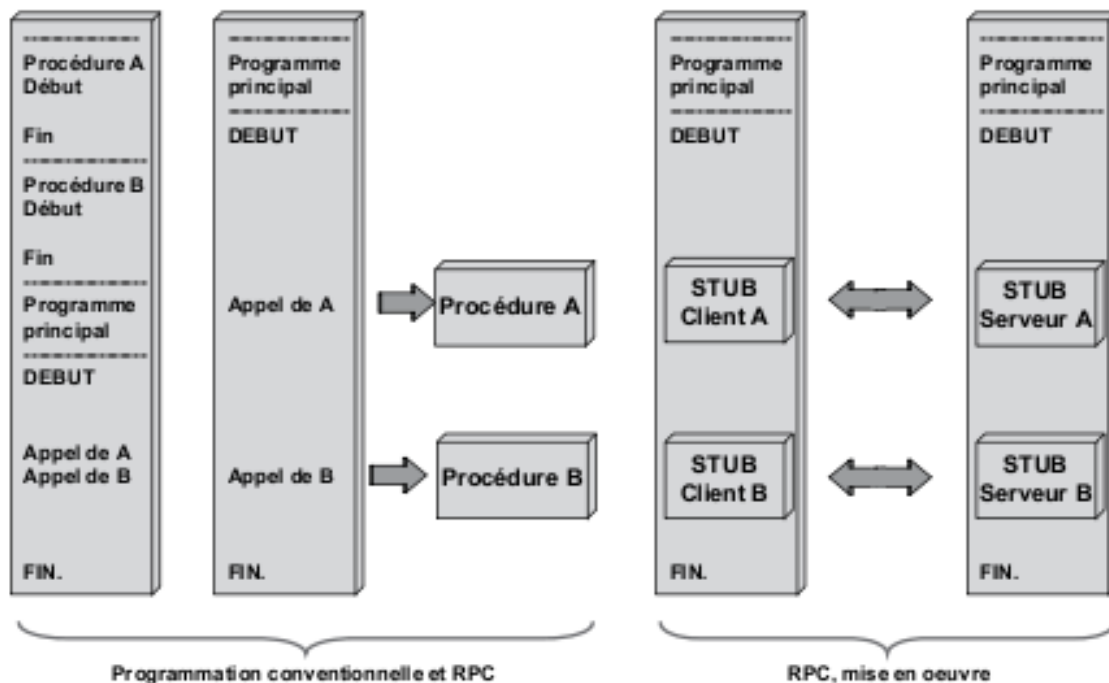


Figure 18.10 La programmation traditionnelle et RPC.

Le RPC ne définit aucune syntaxe de transfert, cette approche facilite les implémentations et garantit l'indépendance totale des systèmes. Cependant, elle nécessite de la part de chaque entité de résoudre le problème de conversion de syntaxe des messages, chaque machine émettant ses données dans son format local.

■ Les sockets

Les *sockets* constituent un élément du *middleware* de bas niveau, ils mettent en relation une application cliente et une application serveur qui ne voient les couches de communication qu'au travers de l'API *socket*. Les *sockets* permettent d'établir un lien en mode connecté ou en mode non connecté à travers le réseau. Ils offrent un ensemble de primitives dont l'enchaînement est représenté figure 18.11.

18.5.3 Internet et le middleware

■ Le Word Wide Web (WWW)

Le WWW constitue la dernière génération des serveurs d'information. Développé par le CERN (Centre européen de la recherche nucléaire), le web est constitué d'un ensemble de serveurs d'information pointant les uns sur les autres par des hyperliens contenus dans les documents consultés. Les principaux logiciels permettant le développement, la mise en œuvre et la consultation d'information organisée en page en mode hypertexte sont les suivants :

- ▶ **HTTP** (*HyperText Transport Protocol*) constitue le protocole de type question/réponse de transfert de fichiers hypertextes entre un serveur web et un client web ;
- ▶ **HTML** (*HyperText Markup Language*) définit le format d'un document et la syntaxe de description des pages, remplacé aujourd'hui par **XML** (*eXtensible Markup Language*) ;
- ▶ **URL** (*Uniform Resource Locators*) définit le format de l'adresse (localisation) d'un document à importer. La syntaxe est la suivante :

Protocole ://Nom du serveur/Répertoire/Nom de fichier.

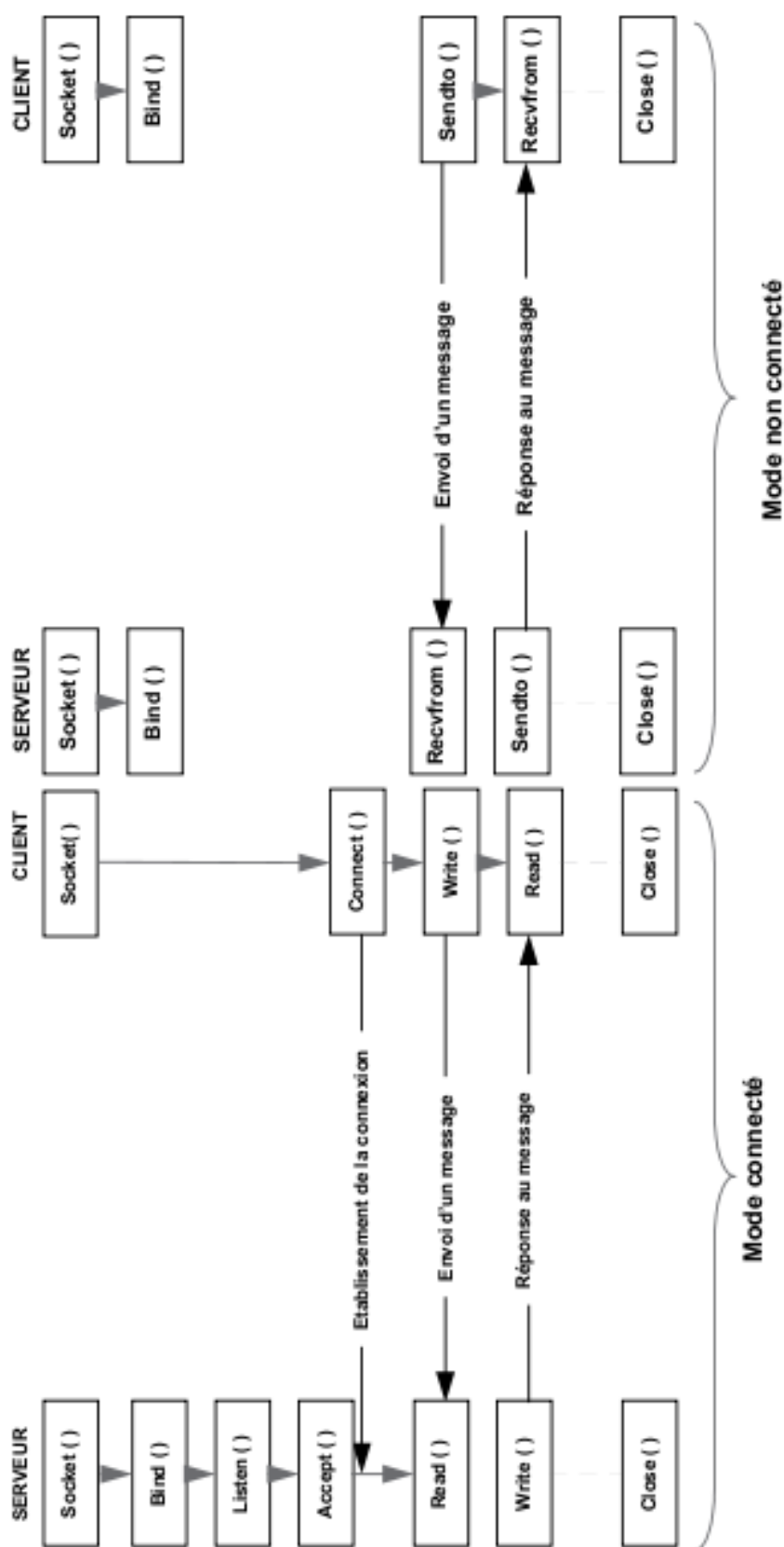


Figure 18.11 L'enchaînement des différentes primitives.

Les URL permettent l'accès depuis le navigateur Web à toutes les formes de documents sur le Web :

- ▶ « http :// » pour les serveurs Web,
- ▶ « ftp :// » pour les serveurs FTP,
- ▶ « gopher :// » pour accéder aux services Gopher.

Un document est dit hypertexte quand il contient des liens vers d'autres documents (hyperliens) et hypermédia s'il pointe vers des documents contenant des images et/ou du son.

■ Le logiciel client

S'il est toujours possible de se connecter à Internet à l'aide des logiciels clients spécifiques à chaque service invoqué, il est plus aisé de s'y connecter au travers d'un logiciel unique dit « navigateur Web ». Ces navigateurs donnent accès aux services d'Internet *via* une interface unique (figure 18.12).

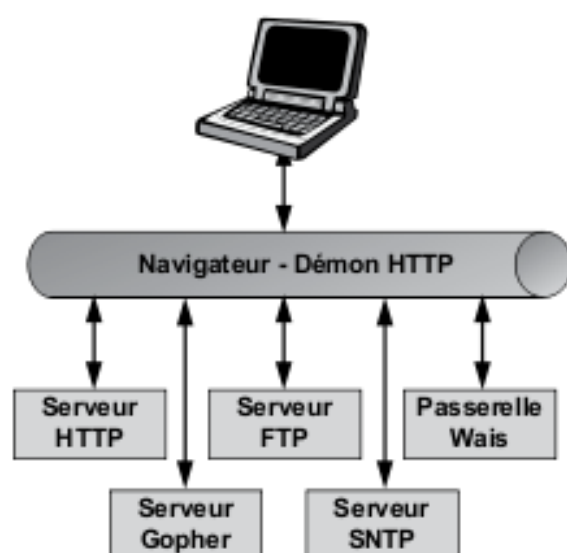


Figure 18.12 Le navigateur, interface aux services d'Internet.

Les navigateurs fonctionnent selon le mode clients/serveur. En principe, le client n'est qu'un terminal passif chargé de l'affichage des pages HTML. Cependant, la puissance de calcul des machines clientes autorise le serveur à invoquer celles-ci pour l'exécution de certaines tâches. L'uti-

lisateur n'a aucun contrôle sur ces petits programmes qui peuvent donc poser des problèmes de sécurité. Ces programmes ou routines dénommés « *applet* » (en français : appliquette !) sont des scripts chargés par le navigateur et interprétés par l'interpréteur « Java » (machine virtuelle). Cette approche assure l'indépendance totale entre le développement et la machine cliente.

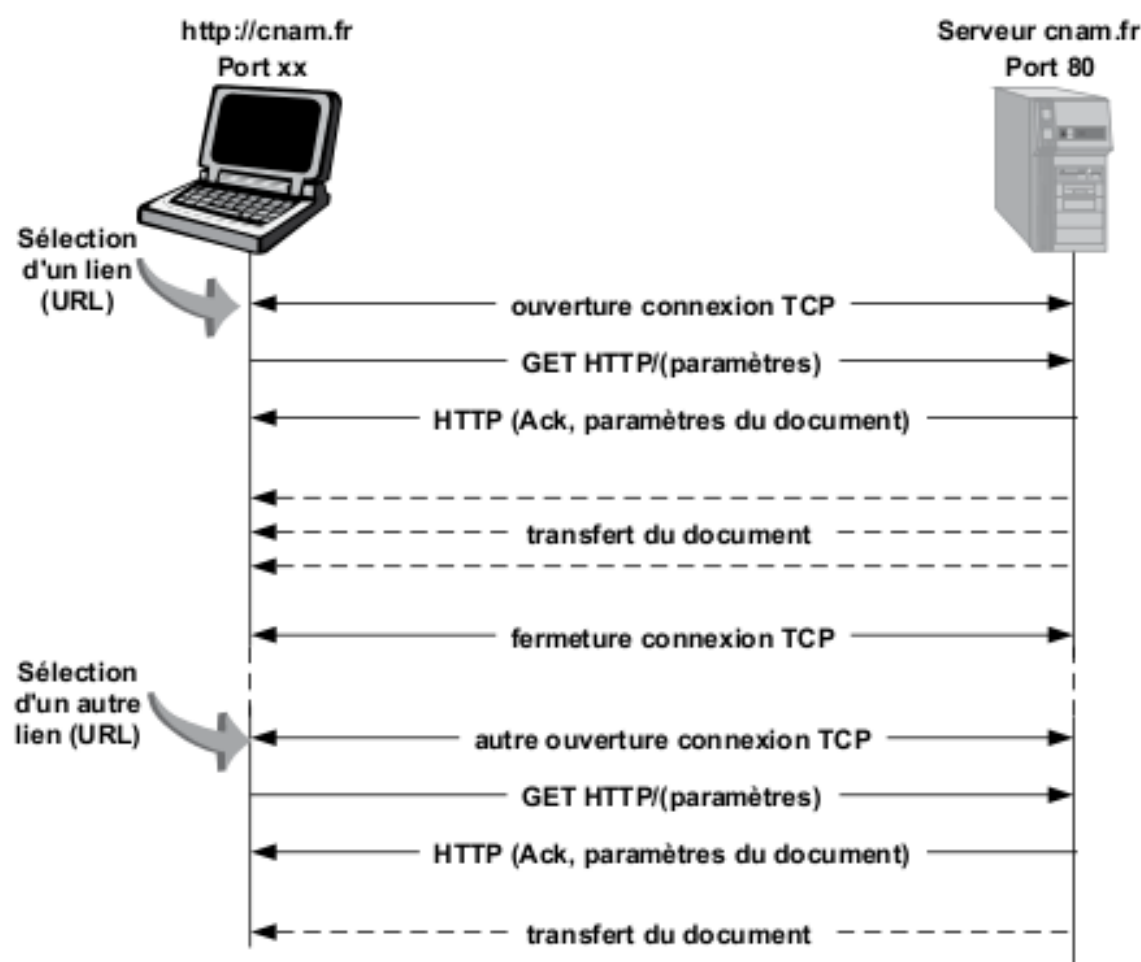


Figure 18.13 Les échanges entre un navigateur et un serveur HTTP.

La figure 18.13 illustre une session d'échange entre un serveur HTTP (port 80¹ par défaut) et un client HTTP.

1 Pour des raisons de sécurité certains serveurs n'utilisent pas le port par défaut. Il faut alors préciser ce dernier lors de l'appel : `http://exemple.fr:8080`, le service est accessible sur le port 8080.

7

Les réseaux locaux

19

Le réseau local

Un réseau local (**LAN**, *Local Area Network*) est un ensemble de moyens autonomes de calcul (micro-ordinateurs, stations de travail ou autres) reliés entre eux pour échanger des informations et/ou partager des ressources matérielles (imprimantes, espace disque...) ou logicielles (programmes, bases de données...).

19.1 Les constituants d'un réseau local

L'architecture OSI organise l'interconnexion de systèmes en mode point à point, alors que les réseaux locaux partagent un support unique en **mode diffusion**. Cette différence d'approche a conduit à redéfinir l'accès au support qui a été implémenté au niveau 2 du modèle de référence de l'ISO. Le niveau liaison (figure 19.1) a été divisé en deux. La sous-couche la plus basse (sous-couche **MAC**, *Medium Access Control*) contrôle l'accès partagé au support et le contrôle d'erreur. La sous-couche supérieure (sous-couche **LLC**, *Logical Link Control* ou Contrôle du lien logique) remplit les fonctions traditionnellement dévolues à la couche liaison (établissement d'un lien logique).

Sur une machine connectée en réseau local, les différentes commandes peuvent être adressées soit au système local soit à un système distant. Il est alors nécessaire de distinguer ces deux types d'appel. Une couche fonctionnelle dite « redirecteur » a pour rôle de diriger les appels vers le système cible. La notion de redirecteur¹ n'est pas définie par l'ISO. Cepen-

¹ Le redirecteur, désignation Microsoft, est couramment appelé *requester* ou *shell* chez Novell.

nant, on peut admettre qu'elle se situe au niveau de la couche présentation. La figure 19.1 illustre cette architecture.

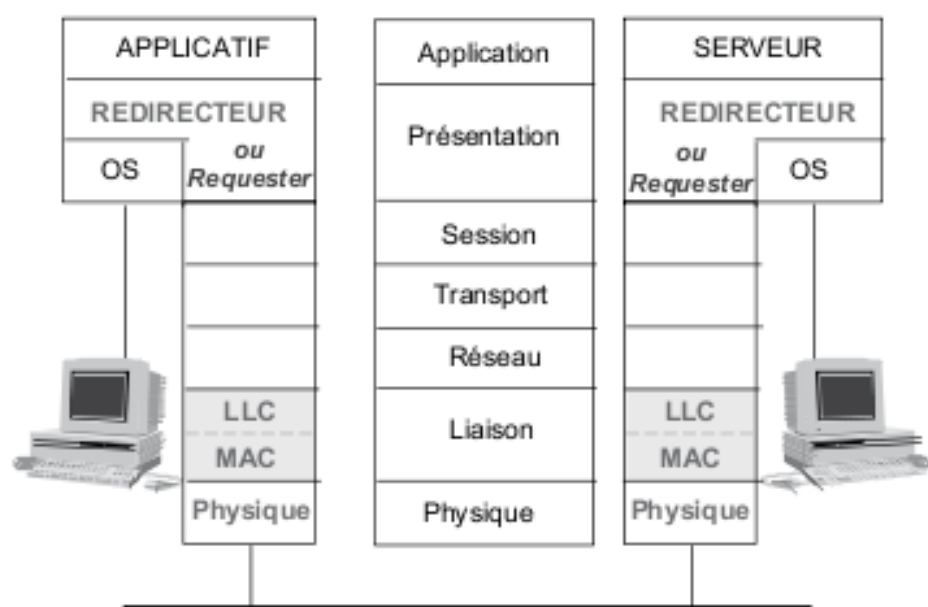


Figure 19.1 L'architecture générale des réseaux locaux.

19.2 Les réseaux locaux et la normalisation

Devant la diversité des besoins et des produits proposés, l'**IEEE** (*Institute of Electrical and Electronic Engineers*) a créé le groupe de travail 802 (février 1980) chargé de définir des standards (Standards 802.x). En 1988, l'ISO a repris la plupart de ces standards pour les normaliser et en faire des normes internationales (série ISO 8802.x).

Aujourd'hui, l'IEEE poursuit, au profit de l'ISO, son travail de normalisation. Le groupe 802 est divisé en sous-groupes de travail, chacun chargé d'une tâche normative dans un domaine spécifique.

19.3 La couche physique

La couche physique spécifie les modes de raccordement (topologie et câblage), les niveaux électriques et le codage des informations émises.

La topologie d'un réseau décrit la manière dont les différents composants du réseau sont reliés. Les réseaux locaux utilisent les topologies de base comme le bus, l'anneau et l'étoile (figure 19.2) ou des combinaisons de celles-ci (étoile de bus, grappe d'étoiles...).

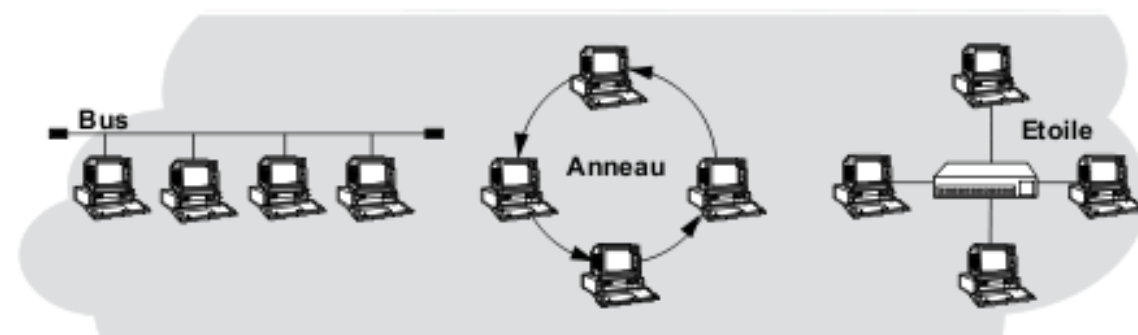


Figure 19.2 Les topologies de base.

Sur un bus, les unités sont au même niveau hiérarchique, les messages sont reçus par l'ensemble des stations (diffusion). Le système n'étant pas hiérarchisé, une station peut accéder au support à tout moment. Ce mode d'accès n'interdit pas à deux stations d'émettre en même temps, les messages sont alors altérés : il y a **collision** ou **contention**. Pour résoudre ce problème, avant d'émettre, la station vérifie qu'aucune autre n'est en émission (écoute du support), cette méthode d'accès est utilisée par les réseaux IEEE 802.3 appelés couramment « **Ethernet**¹ ».

L'anneau, cas particulier d'une liaison multipoint, implique une circulation unidirectionnelle des messages. Dans ce type de topologie, le droit d'émettre (jeton) est transmis à la station qui suit physiquement celle qui le détient (jeton non adressé). Cette méthode d'accès est mise en œuvre dans le réseau IEEE 802.5 ou **Token Ring**, aujourd'hui obsolète.

19.4 La sous-couche MAC

La sous-couche **MAC** (*Medium Access Control*) gère l'accès au support physique, elle règle les problèmes d'adressage (adresse MAC) et effectue un contrôle d'erreur (**FCS**, *Frame Check Sequence*).

¹ *Ethernet* est un nom de marque déposé par Xerox. Ce nom est passé dans le langage courant et désigne les réseaux de type CSMA/CD.

19.4.1 Les méthodes d'accès

Ce sont les méthodes d'accès qui distinguent les différents types de réseau et déterminent leurs performances dans tel ou tel environnement. Deux méthodes ont dominé le monde des réseaux locaux : les méthodes aléatoires ou à contention, mises en œuvre dans les réseaux de type « Ethernet », et les méthodes à réservation fondées sur le passage du droit d'émettre (jeton) dont le Token Ring a été l'implémentation la plus connue.

19.4.2 L'adressage MAC

■ Généralités

L'adresse MAC (adresse physique) désigne de manière unique une station sur le réseau. À des fins de facilité d'administration, elle est gravée dans l'adaptateur réseau (**NIC**, *Network Interface Card*) par le fabricant de l'adaptateur. L'IEEE est chargée de la gestion de ces adresses, elle en garantit l'unicité en attribuant un identifiant différent à chaque fabricant d'interface réseau (figure 19.3).

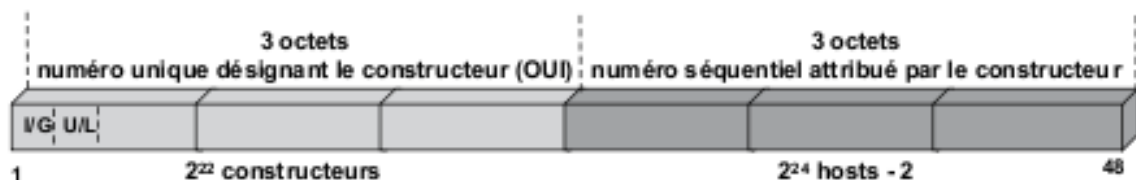


Figure 19.3 L'adressage IEEE.

Le premier bit (bit I/G) distingue une adresse individuelle ou *unicast* ($I = 0$) d'un adressage de groupe (*multicast* ou *broadcast*, $I = 1$). Le bit suivant (bit U/L) détermine si l'adresse qui suit est universelle : adressage IEEE ($U = 0$) ou local ($L = 1$). Dans ce dernier cas, c'est l'administrateur de réseau qui gère l'espace d'adressage et garantit l'unicité d'adressage.

Dans l'adressage universel, les 22 bits suivants désignent le constructeur ou le revendeur de l'adaptateur réseau. L'IEEE attribue à chaque constructeur un ou plusieurs numéros qui l'identifient (**OUI**, *Organization Unit Identifier*). Les 24 bits suivants appartiennent à une série séquentielle et sont inscrits sous la responsabilité du fabricant (**SN**, *Serial Number*) dans l'adaptateur. L'adressage IEEE est un adressage à plat, il distingue, sur le

réseau, une machine parmi les autres, mais ne permet pas d'en déterminer la position géographique.

■ Les différentes adresses MAC

□ L'adresse individuelle ou *unicast*

L'adresse *unicast* est utilisée pour les échanges entre stations. L'adresse *unicast* est gravée dans la carte lors de sa fabrication, mais l'administrateur du réseau peut lui substituer une adresse *d'unicast* locale (bit L à 1).

□ L'adresse de diffusion généralisée ou *broadcast*

Une adresse de *broadcast* est une adresse de diffusion générale. Tous les bits de l'adresse MAC sont positionnés à 1 (FF-FF-FF-FF-FF-FF).

□ L'adresse de diffusion restreinte ou *multicast*

Une adresse de *multicast* ou de groupe (bit G = 1) désigne un ensemble de stations. Les applications fournissent à la station (couche MAC) la liste des adresses de groupe auxquelles elle doit répondre (abonnement). Ces adresses sont utilisées par exemple pour la diffusion vidéo. Des plages d'adresses *multicast* ont été définies pour permettre l'encapsulation d'adresses IP *multicast*, cette plage s'étend de :

01-00-5E-00-00-00 à 01-00-5E-7F-FF-FF (RFC 1112)

Pour autoriser une diffusion sélective, toutes les adresses MAC du groupe de diffusion doivent être identiques. À cet effet, l'adresse MAC *multicast* de chaque station du groupe est construite à partir de l'adresse IP *multicast* du groupe. La figure 19.4 illustre la détermination d'une adresse *multicast* IEEE (adresse MAC) à partir d'une adresse IP *multicast* (classe D).

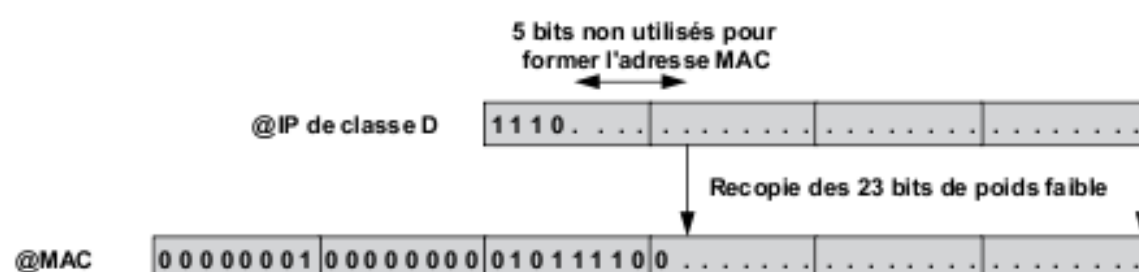


Figure 19.4 Détermination de l'adresse @MAC multicast.

19.5 La couche liaison (LLC)

19.5.1 Généralités

La sous-couche LLC¹ (*Logical Link Control*) assure un service comparable à celui offert par la couche liaison du modèle de référence. Elle masque à la couche supérieure le type de réseau utilisé (Ethernet, Token Ring...). La couche LLC offre, selon les besoins, trois types de services : LLC1, LLC2 et LLC3.

19.5.2 Le service LLC de type 1

Le service LLC1 est un service en mode datagramme. La couche MAC réalise un contrôle d'erreur, mais n'effectue pas de reprise sur erreur, toute trame erronée est purement et simplement éliminée. C'est le service le plus simple et pratiquement le seul utilisé dans les réseaux locaux.

LLC1 utilise la trame de type UI (*Unnumbered Information*, champ de contrôle à 0x03) représentée figure 19.5.

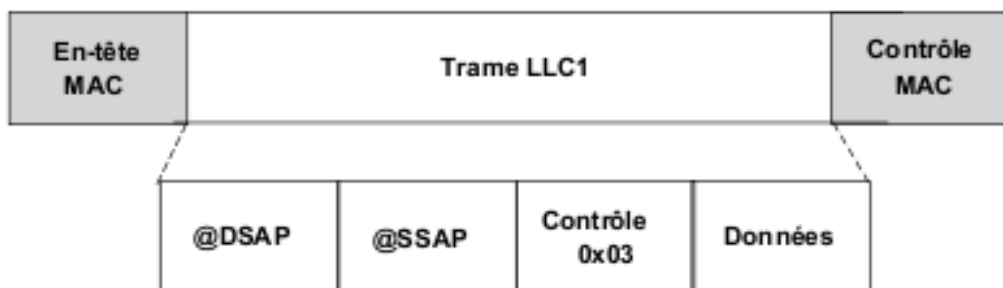


Figure 19.5 Le format de la trame LLC1.

1 La couche LLC n'est en principe plus utilisée dans les réseaux locaux de type Ethernet pour le transfert de données. Mais l'encapsulation LLC (LLC1) en conjugaison avec l'encapsulation SNAP est utilisée par la plupart des utilitaires (*spanning tree...*) et dans de nombreux autres environnements (ATM...).

19.5.3 La sous-couche SNAP (*SubNetwork Access Protocol*)

Les 7 bits des champs SAP de la trame LLC ne permettent pas d'identifier tous les protocoles de niveau supérieur (espace de numérotation insuffisant). L'encapsulation **SNAP** (*SubNetwork Access Protocol*) introduit un champ d'identification supplémentaire **PIH** (*Protocol Identifier Header*). La figure 19.6 représente l'encapsulation SNAP pour une trame MAC IEEE 802.3 (Ethernet aussi dans le langage courant). L'encapsulation SNAP est identifiée par les champs DSAP et SSAP qui contiennent la valeur 0xAA.

Le **PIH** (*Protocol Identifier Header*) est divisé en deux champs : le champ **OUI** (*Organization Unit Identifier*) et le champ **PID** (*Protocol Identifier*). La valeur « 0 » du champ OUI indique que le champ PID est codé de la même façon que le champ Ethertype de la trame Ethernet (voir figure 20.3).

L'encapsulation SNAP est utilisée par les utilitaires réseaux de la famille 802.1 et dans les protocoles haut débit (Frame Relay, RFC 1490 et ATM, RFC 1483).

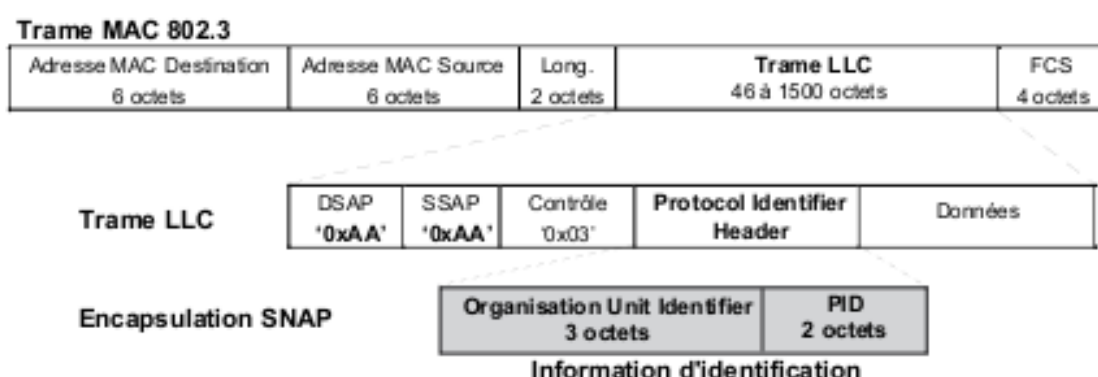


Figure 19.6 L'encapsulation SNAP.

20 Les réseaux Ethernet

20.1 Présentation

Passant de 3 Mbit/s pour la version expérimentale à 40 Gbit/s voire 100 Gbit/s, en évoluant d'un environnement à contention vers un environnement commuté, en migrant du LAN Ethernet aux MAN voire WAN, en devenant le support de tous types d'application y compris le multimédia, Ethernet a su, en permanence, s'adapter à l'évolution des besoins. Le tableau 20.1 fournit une synthèse de ces évolutions.

Tableau 20.1 La synthèse des évolutions des réseaux Ethernet.

Année	Nom	Description succincte
1976	Alto Aloha Network	Réseau expérimental à 2,94 Mbit/s.
1980	Ethernet DIX	Pré-version de l'Ethernet sur coaxial épais à 10 Mbit/s.
1982	Ethernet v2	Version de base des réseaux Ethernet, repris par l'IEEE sous le nom de IEEE 802.3 10Base5 en 1985, puis en 1989 par l'ISO (ISO 8802.3).
1988	IEEE 802.3a ou 10Base2	Extension de la norme sur coaxial fin.
	IEEE 802.b ou 10BROAD36	Version en large bande.
	IEEE 802.c	Définition des répéteurs pour prolonger le réseau.
	IEEE 802.e ou 1Base5	Reprise des spécifications du Starlan, version d'Ethernet sur paires torsadées.
1989	IEEE 802.d Ethernet FOIRL	Ethernet 10 Mbit/s sur fibre optique. (FOIRL, Fiber Optic Inter-Repeater Link).
1990	IEEE 802.3i 10BaseT	Extension du réseau Starlan à 10 Mbit/s.
	IEEE 802.3h	Définition de l'administration (Layer Management).

Année	Nom	Description succincte
1993	IEEE 802.3j 10BaseF	Redéfinition de FOIRL, apport d'une configuration en étoile avec répéteurs multi-ports.
	IEEE 802.3q	Agrégation de VLAN.
1995	IEEE 802.3u Fast Ethernet	Définition des versions à 100 Mbit/s : 100BaseTx, 100BaseT4, 100BaseFx.
1998	IEEE 802.3x	Introduction du fonctionnement en full duplex et du contrôle de flux.
	IEEE 802.3y	Définition de la version à 100 Mbit/s 100BaseT2.
	IEEE 802.3z GigabitEthernet	Définition d'une version à 1 Gbit/s sur fibre optique (1 000 BaseLX, BaseSX) et paires torsadées blindées (1 000 BaseCX).
	IEEE 802.3ab 1000BaseT	Définition d'une version à 1 000 Mbit/s sur les infrastructures existantes (câble catégorie 5).
	IEEE 802.3ac	Définition du marquage des VLAN (VLAN TAG).
2000	IEEE 802.3ad	Définition de l'agrégation de liens (Link Aggregation) Reclassé en 802.1 en 2008.
2002	IEEE 802.3ae	Version à 10 Gbit/s (10GBASE-F).
2003	IEEE 802.3af	Auto-alimentation des terminaux (PoE).
2004	IEEE 802.3ah	Ethernet in the First Mile.
2006	IEEE 802.3an	Ethernet 10 Gbit/s sur paires torsadées (10GBASE-T).
2008	IEEE 802.3at	Évolution de PoE (Power over Ethernet).
	IEEE 802.3ax	Déplacement de Link Aggregation de 802.3 vers 802.1.
2009	IEEE 802.3av	10 Gbit/s PHY EPON.
	IEEE 802.3ba	Ethernet à 40 et 100 Gbit/s.

Ethernet a pour origine la méthode d'accès Aloha du réseau radio de l'université d'Hawaï (1970). Une station qui avait un message à transmettre l'émettait sans se préoccuper des autres stations. Si le message était pollué par une autre émission (collision), il était retransmis. Cette méthode, très simple à implémenter, est d'autant moins efficace que le nombre de stations

augmente. La méthode dite **CSMA/CD** (*Carrier Sense Multiple Access, Collision Detection*) dérive de cette approche. D'abord baptisé Alto Aloha, l'Ethernet (réseau dans l'Ether) a été développé chez Xerox par Robert Metcalfe dans les laboratoires d'Alto (**Alto Aloha Network**) en 1973. Associant Digital, Intel et Xerox (**DIX**), Bob Metcalfe fit évoluer sa version vers 10 Mbit/s (Ethernet DIX, 1980). Les spécifications définitives (Ethernet v2, 1982) ont été reprises par l'IEEE pour donner naissance aux spécifications IEEE 802.3 10Base5 (1985), puis par l'ISO (ISO 8802.3) en 1989. Aujourd'hui, cette technique domine largement le marché et fait un retour aux sources avec le développement de la norme 802.11 (réseaux sans fil).

20.2 Caractéristiques des réseaux Ethernet

20.2.1 Le principe du CSMA/CD

Le principe de base du CSMA/CD repose sur la **diffusion** des messages à l'ensemble des stations du réseau (réseau à diffusion). La figure 20.1 illustre la mise en œuvre du CSMA/CD dans les réseaux 802.3 et Ethernet.

La station *A* diffuse son message (t_0 à t_3). En t_1 , la station *B*, avant d'émettre, se met à l'écoute. Le support est occupé, elle diffère son émission, mais reste à l'écoute (attente active). De même *C*, en t_2 , se porte à l'écoute et retarde son émission. En t_3 , *A* cesse d'émettre, *B* et *C* détectent le silence, ils émettent en même temps. En t_4 , chacune des stations détecte que son message est altéré, la collision est détectée. *B* et *C* cessent leur émission et déclenchent une temporisation aléatoire. En t_5 , le *timer* de *B* arrive à échéance. Le canal étant libre, *B* émet. En t_6 , *C* détecte le support occupé, il diffère son émission jusqu'au temps t_7 .

20.2.2 La fenêtre de collision

La fenêtre de collision correspond au temps minimal pendant lequel une station doit émettre pour détecter la collision la plus tardive que son message est susceptible de subir. Ce temps minimal d'émission s'appelle fenêtre de collision, *Time Slot* ou encore tranche canal. C'est la période de vulnérabilité d'un message, c'est-à-dire l'intervalle de temps pendant lequel une collision peut se produire.

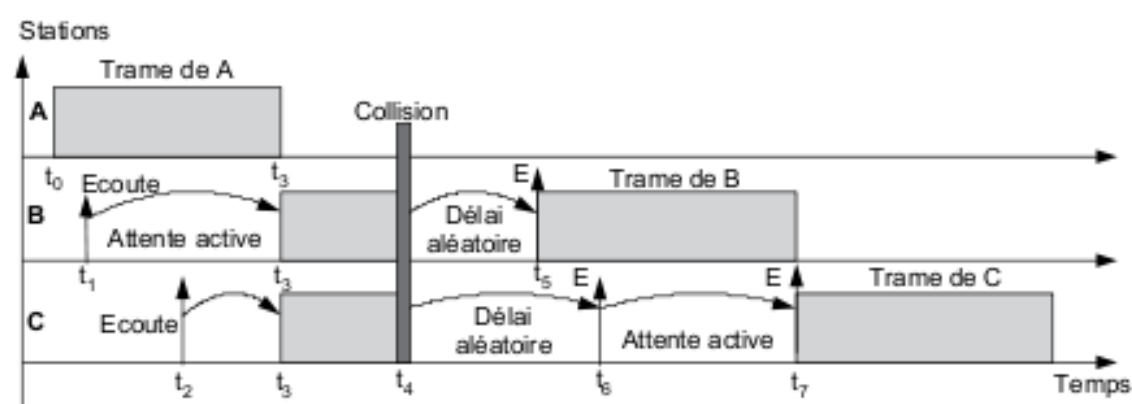


Figure 20.1 Principe du CSMA/CD.

Ce temps minimal d'émission correspond à deux fois le temps de propagation d'une trame sur la plus grande distance du réseau, fixé à $51,2 \mu\text{s}$, ce temps correspond, pour un débit de 10 Mbit/s , à l'émission de 512 bits, soit 64 octets. Si, le message à transmettre est de longueur inférieure à 64 octets, une séquence de bourrage (Padding) est insérée derrière les données utiles. Un pointeur de longueur des données utiles, sur deux octets, permet au récepteur d'extraire les données utiles de l'ensemble des données du champ information (données utiles + bourrage).

La signification du champ « Longueur des données utiles » diffère selon que la trame est au format Ethernet v2 ou au format IEEE 802.3, rendant ces deux types de réseaux CSMA/CD incompatibles (figure 20.2).

20.2.3 Le format des trames Ethernet/IEEE 802.3

La trame IEEE 802.3 est représentée en figure 20.2. Un préambule de 7 octets permet la synchronisation bit. La synchronisation caractère est assurée par le fanion de début de trame (**SFD**, *Start Frame Delimitator*), les deux bits à « 1 » marque le début de la trame. Les champs adresses (2 ou 6 octets) contiennent les adresses MAC destination et source. Un pointeur, sur 2 octets, indique la longueur utile du champ données. Le champ données est suivi d'un contrôle d'erreur de type CRC : le **FCS** (*Frame Check Sequence*) sur 4 octets.

La trame Ethernet DIX (Dec, Intel, Xerox) diffère de la trame IEEE 802.3. Le champ longueur de données de la trame IEEE 802.3 se substitue au champ d'identification du protocole supérieur (EtherType) de la trame

Ethernet DIX ou Ethernet v2. Le champ Type/Longueur permet de distinguer les deux versions. Sur un réseau de type Ethernet, en principe, les données s'échangent au format Ethernet alors que les protocoles de gestion utilisent le format 802.3.

TRAME 802.3

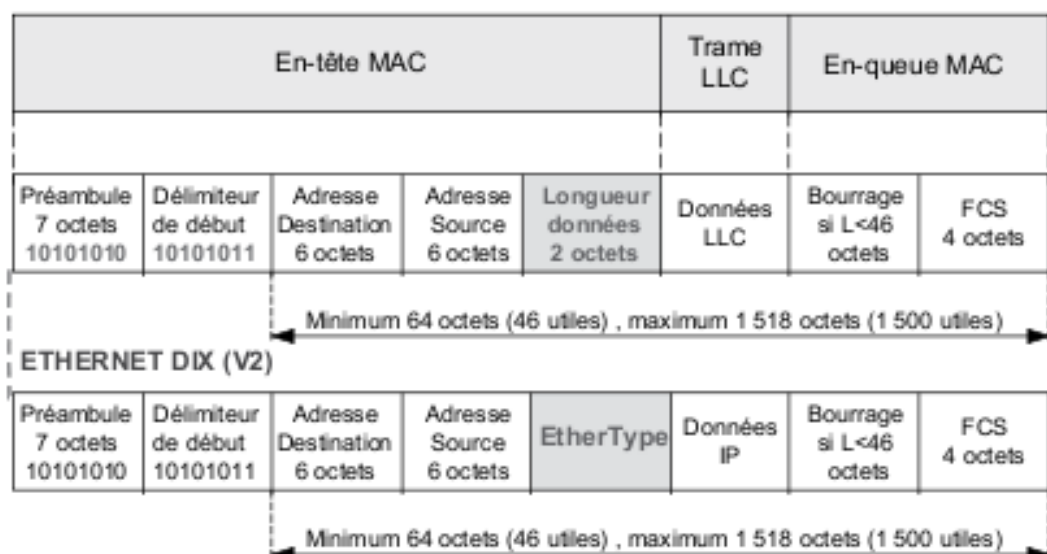


Figure 20.2 Le format de la trame Ethernet/IEEE 802.3.

La couche MAC réalise un contrôle d'intégrité. L'en-queue de la trame MAC contient le champ de contrôle d'erreur par CRC sur 32 bits (FCS, *Frame Control Sequence*). Le polynôme générateur, identique pour tous les types de réseaux locaux normalisés par l'IEEE, est :

$$x^{32} + x^{26} + x^{23} + x^{22} + x^{16} + x^{12} + x^{10} + x^8 + x^7 + x^5 + x^4 + x^2 + 1$$

La couche MAC rejette toute trame erronée mais n'effectue aucune reprise sur erreur. Cette tâche est reportée sur les couches supérieures.

20.3 Les différentes versions d'Ethernet

20.3.1 L'Ethernet épais, IEEE 802.3 10Base5

Les appellations IEEE désignent d'abord le sous-comité (802.3), le type de modulation (bande de base ou large bande : *Broad band*) et le diamètre

du réseau ou le type de support. La version 10Base5, (10 Mbit/s en bande de base sur câble coaxial d'une longueur maximale par segment de 500 m) utilise un codage Manchester. La figure 20.3 illustre cette version d'origine de l'Ethernet.

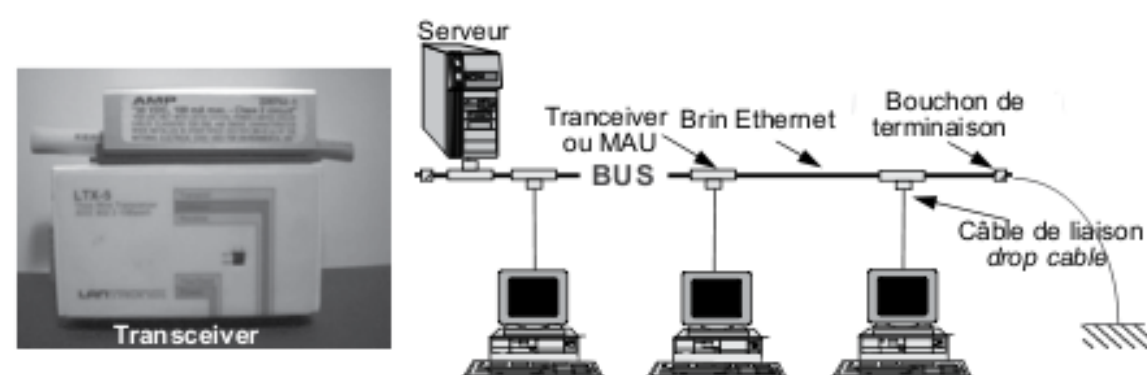


Figure 20.3 Le réseau Ethernet jaune.

Cette version d'Ethernet n'est plus que très rarement utilisée. Elle subsiste dans les environnements compromis (rayonnement électromagnétique) ou lorsque l'on veut garantir la confidentialité des échanges (pas de rayonnement du câble coaxial) qui n'ont pas encore migré vers des versions en fibre optique.

20.3.2 L'Ethernet fin, IEEE 802.3 10Base2

Compte tenu des difficultés de câblage de la version 10Base5, une version économique a été définie sur du câble coaxial fin (*Thin Ethernet*). Dans cette version, les fonctions du *transceiver* sont remplies par la carte transporteur (MAU intégré à la carte). De ce fait, le bus coaxial est connecté directement sur la carte par l'intermédiaire d'un T vissé BNC (*Barrel Neck Connector*). La longueur maximale d'un segment est de 185 m et chaque segment peut accueillir un maximum de 30 stations. La figure 20.4 présente cette version du réseau Ethernet.

Cette architecture physique de réseau a été très utilisée pour réaliser de petits réseaux d'une dizaine de machines.

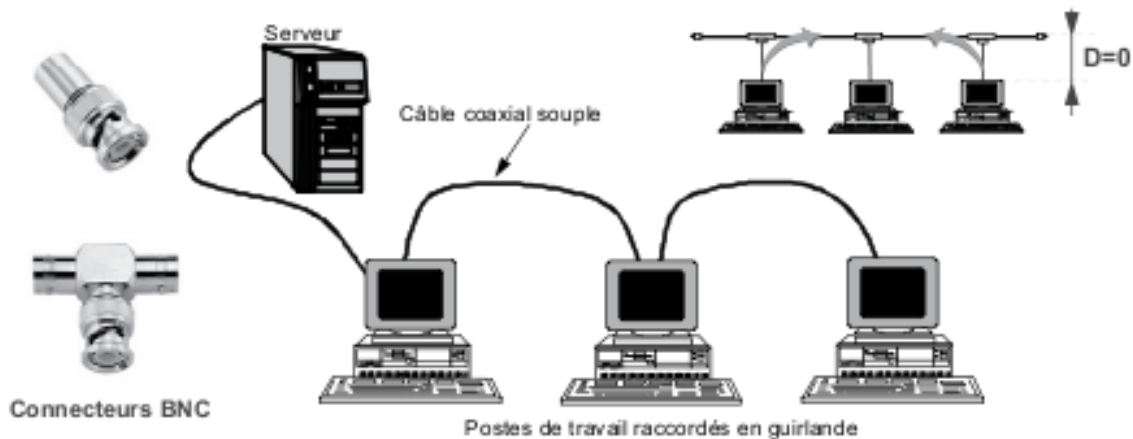


Figure 20.4 L'Ethernet fin.

20.3.3 L'Ethernet sur paires torsadées

Compte tenu des difficultés du câblage coaxial, AT&T a imaginé réutiliser le câblage téléphonique préexistant dans les immeubles de bureaux pour la réalisation de réseau. Le réseau devait alors passer d'une topologie physique et logique en bus à une topologie physique en étoile, assurer la diffusion des messages (topologie logique en bus) et la détection des collisions. La solution adoptée par AT&T consiste simplement à émuler un bus dans un boîtier : le **hub**. Le hub est chargé d'une part de concentrer les connexions et, d'autre part, d'assurer la diffusion des messages (figure 20.5).

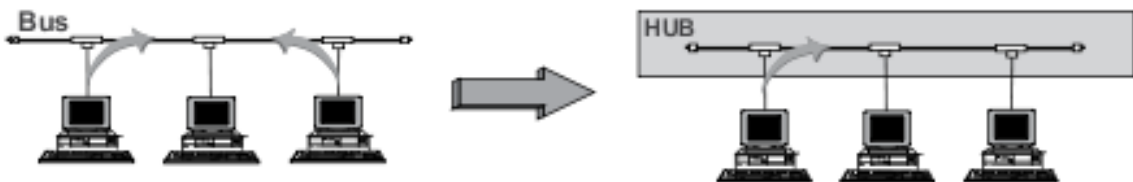


Figure 20.5 Le passage du bus à l'étoile.

La liaison hub/station est réalisée en paires torsadées (1 paire émission, 1 paire réception), ce qui impose deux contraintes : l'une de débit, l'autre de distance. Cette architecture (802.3 1Base5 ou Starlan, limitée à 1 Mbit/s), complètement obsolète, est à l'origine de toutes les évolutions d'Ethernet.

20.3.4 L'Ethernet, 802.3 10BaseT

La version 10BaseT reprend les principes architecturaux du réseau Starlan, la topologie physique reste une étoile hiérarchisée (figure 20.6).

Un signal particulier, le *link status* (état de la ligne), permet, par la visualisation de diodes LED (*Light Emitting Diode*), de contrôler la continuité du lien entre le hub et la station (*Link integrity test function*). En l'absence d'émission, le hub et la carte réseau émettent, toutes les 16 ± 8 ms, une impulsion de test (*TP_IDL*, *Twisted Pair Idle Signal*), pour contrôler l'intégrité du lien (*LTP*, *Link Test Pulses*). En l'absence de signal (données utilisateur ou signal LTP), le hub et la carte considèrent le lien défectueux, le voyant *link status* est alors positionné à « OFF » et le port du hub est inhibé.

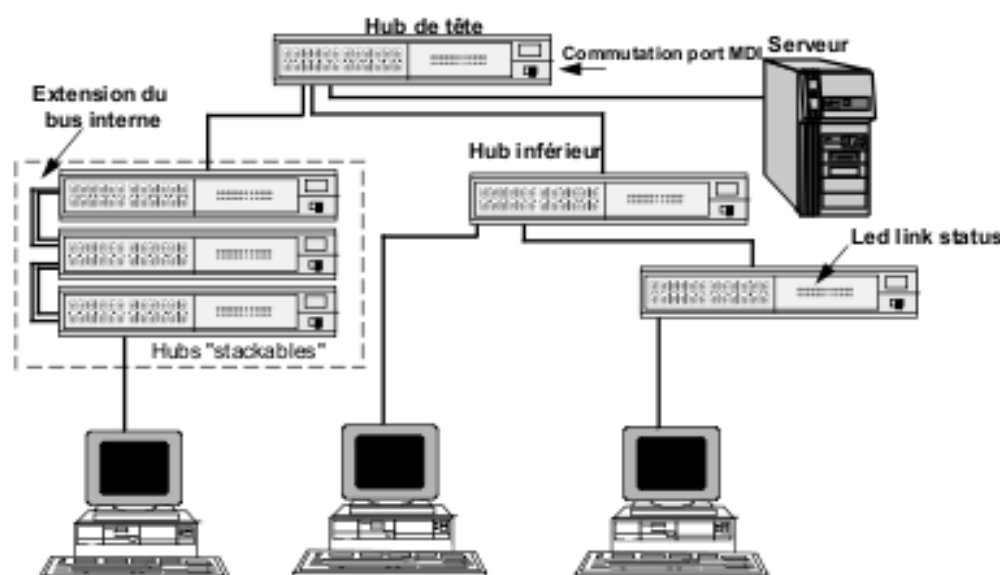


Figure 20.6 – L'architecture du réseau 10BaseT.

20.3.5 L'Ethernet à 100 Mbit/s

■ Généralités

L'Ethernet 100 Mbit/s sur paires torsadées, ou *Fast Ethernet*, est l'évolution naturelle de la version 10BaseT. La compatibilité avec la version 10BaseT est assurée par la reprise du protocole CSMA/CD et le maintien des tailles de trame (64 octets au minimum et 1 518 au maximum). De ce fait, la fenêtre de collision ou tranche canal est réduite à $5,12 \mu\text{s}$ (512 bits). La réduction de la fenêtre de collision par un facteur de 10 (de $51,2 \mu\text{s}$ à $5,12 \mu\text{s}$) induit de fortes contraintes sur le temps de propagation du signal et, par conséquent, sur la distance maximale entre les deux stations les plus éloignées du réseau.

■ La mixité et l'auto-négociation

Pour permettre l'évolution des réseaux en douceur, les ports ont la faculté de s'autoconfigurer à 10 ou 100 Mbit/s. À cet effet, les équipements 100 Mbit/s remplacent le signal de *link status* du 10BaseT par un mot de 16 bits (LCW, *Link Code Word*) décrivant les caractéristiques de l'équipement. À défaut de détection par l'utilisation du mot LCW, le système peut s'autoconfigurer par reconnaissance du type de codage des signaux (4B/5B pour le 100BaseTX et Manchester pour le 10BaseT).

20.3.6 Le Gigabit Ethernet

En portant le débit à 1 Gbit/s, la version d'Ethernet dite Gigabit Ethernet (**GbE**) introduit selon le type de concentrateur (hub) utilisé et, en environnement partagé, deux modes de fonctionnement : le mode dit unidirectionnel ou *half duplex* (la détection de collision est réalisée par la carte réseau du poste de travail) et le mode bidirectionnel ou *full duplex* (la détection de collision est assurée par le hub). Dans les deux cas, une seule station peut être raccordée à un port. Cependant, le GbE s'est essentiellement développé en mode commuté. L'architecture générale du Gigabit est représentée figure 20.7.

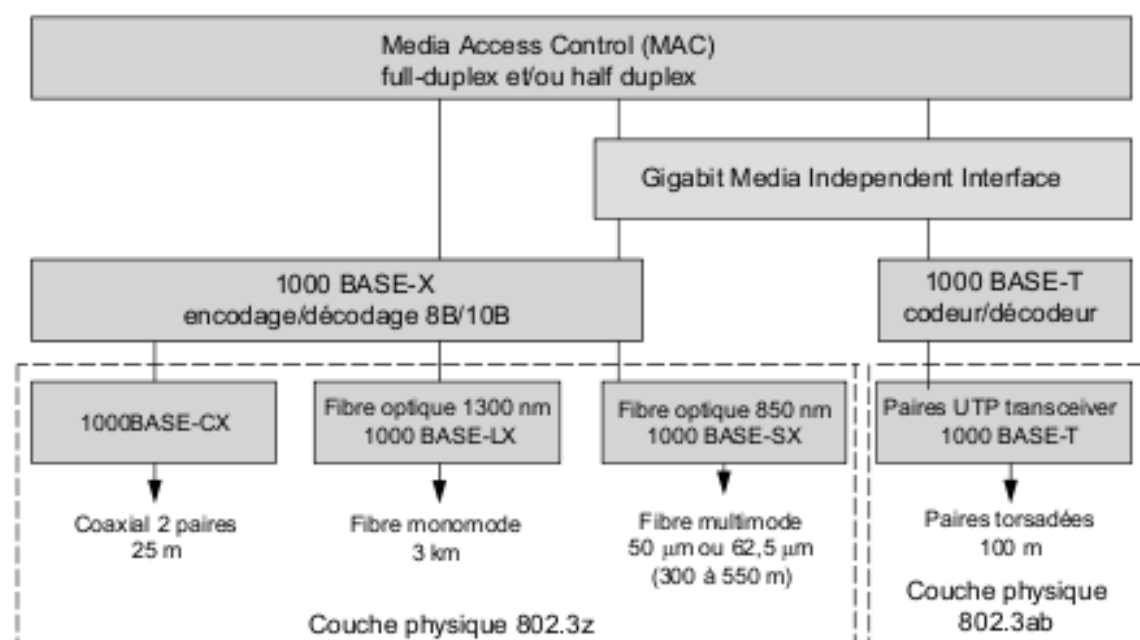


Figure 20.7 L'architecture du Gigabit Ethernet.

20.3.7 Le 10 Gigabit Ethernet

Avec la publication IEEE 802.3ae, Ethernet multiplie encore le débit par 10. Le standard Ethernet 10 Gigabit (10 GE) est essentiellement destiné à une exploitation dans les dorsales de réseaux. Le Gigabit Ethernet et le 10 Gigabit sortent du LAN pour offrir des solutions d'encapsulation de niveau 2 aussi bien dans les MAN, les WAN et même sur la boucle locale (*Ethernet First Mile* ou IEEE 802.3ah).

Le 10 GE conserve la compatibilité avec les versions précédentes. Définit uniquement en mode *full duplex*, le 10 GE offre des interfaces LAN, MAN et WAN sur une distance de 40 km. En s'interfaçant avec des interfaces Sonet/SDH, le 10GE est susceptible de réaliser des liaisons de plusieurs milliers de kilomètres. La figure 20.8 illustre les différentes architectures du 10 GE.

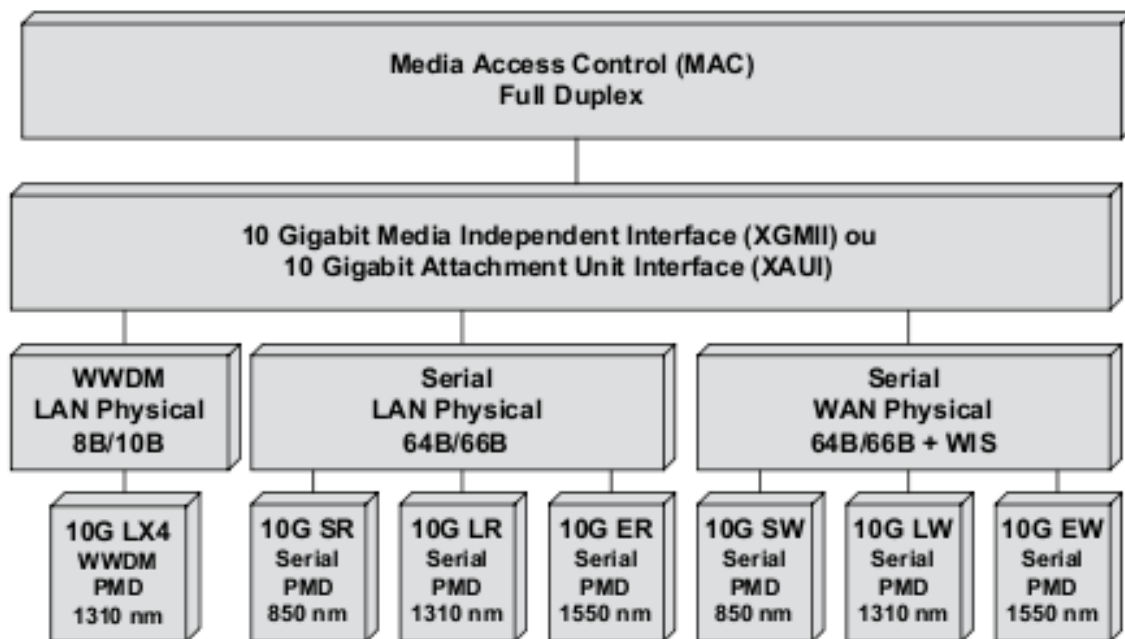


Figure 20.8 L'architecture générale du 10 GE.

21

La commutation dans les LAN – Les réseaux virtuels ou VLAN

21.1 Principe de la commutation dans les LAN

Issue de la téléphonie et des réseaux grande distance (WAN), puis mis en œuvre dans le monde Ethernet (*Switched Ethernet*) pour résoudre les problèmes d'effondrement des réseaux (collisions) et garantir une certaine bande passante, la technique de commutation est aujourd'hui largement utilisée pour réaliser tout type de réseaux. Traditionnellement, la commutation consiste, en fonction d'un identifiant (label), à mettre en relation directe un port d'entrée avec un port de sortie, la relation étant établie préalablement à toute communication par un protocole de signalisation (table de commutation).

La commutation dans les réseaux locaux n'ouvre pas explicitement un circuit virtuel. À l'instar des ponts dont ils ne sont qu'une évolution (ponts multiports), les commutateurs par auto-apprentissage établissent dynamiquement (commutation dynamique), par écoute du trafic, une table de localisation ou d'acheminement (figure 21.1).

Pour construire sa table d'acheminement (FDB, *Forwarding Data Base*), le commutateur examine le trafic reçu par chacun de ses ports et associe au port l'adresse MAC source de la trame reçue. Le commutateur apprend ainsi la localisation géographique des stations. À réception d'une trame, le commutateur consulte la table d'acheminement (table de commutation) et achemine la trame reçue sur le seul port où est localisé le destina-

taire autorisant ainsi un acheminement en simultané de plusieurs trames (figure 21.1).

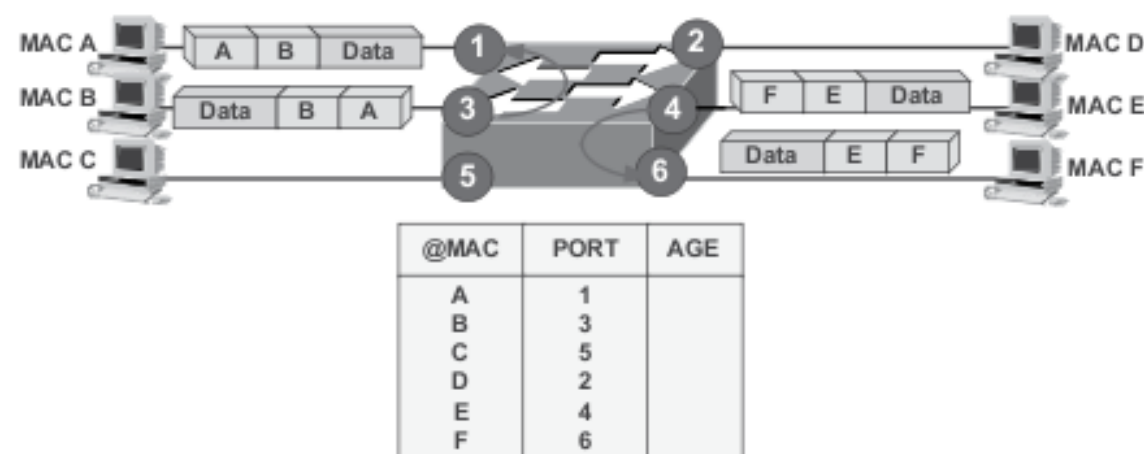


Figure 21.1 La commutation autorise le parallélisme des communications.

Les trames à destination d'une adresse non inscrite dans la table sont répétées sur tous les ports (diffusion), sauf le port de réception. Les tables ne pouvant posséder autant d'entrées que de stations connectées, périodiquement, les adresses les plus anciennes sont effacées. À cet effet, à chaque entrée de la table, est associé un temporisateur réinitialisé à chaque réception d'une trame de même origine.

21.2 Ethernet *full duplex*

Chaque station est reliée à un port, c'est une liaison en mode point à point et par conséquent, les risques de collisions sont inexistantes. L'adaptateur peut alors émettre et recevoir en même temps des messages différents, l'échange est *full duplex*. La technologie *full duplex* (FDSE, *Full Duplex Switched Ethernet*) permet de doubler la bande passante d'un réseau local. Initialement réservée aux liens inter-commutateurs, la technologie *full duplex* est aujourd'hui supportée par la plupart des adaptateurs.

21.3 Principes généraux des VLAN

Avec l'accroissement des stations sur un même réseau, les messages de diffusion (ARP, annonces de service...) occupent une part de plus en plus importante de la bande passante. En définissant, indépendamment de la situation géographique des systèmes, des domaines de diffusion (domaine de *broadcast*), les VLAN autorisent une répartition et un partage optimal des ressources de l'entreprise. Application directe de la commutation statique, les VLAN associent à un port un identifiant. Ne peuvent communiquer que les machines raccordées à des ports de même identifiant. Ainsi, sur le commutateur de la figure 21.2 deux VLAN sont déclarés. La communication entre stations n'est possible qu'entre les stations A, C et D d'une part et les stations B, E et F d'autre part. Il en est de même pour les *broadcasts* qui ne sont diffusés qu'au sein de leur VLAN respectif (domaine de diffusion).

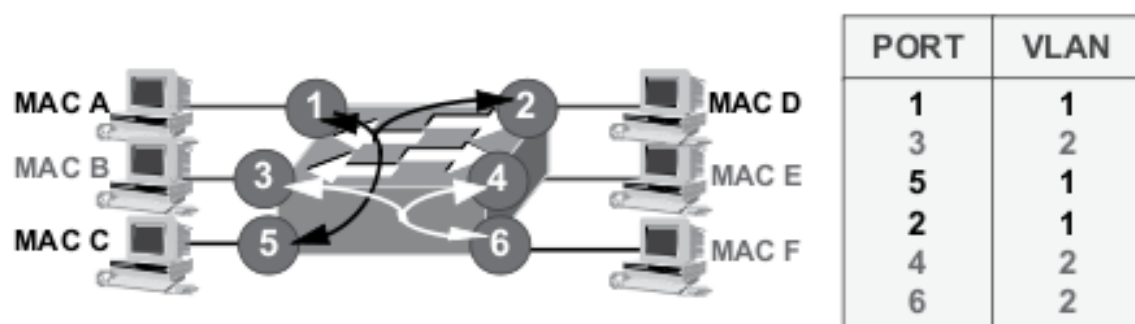


Figure 21.2 Principe des VLAN.

La communication n'est autorisée qu'entre machines d'un même VLAN. Les communications inter-VLAN doivent transiter par un routeur (figure 21.3). Ainsi, les réseaux virtuels permettent de réaliser des réseaux axés sur l'organisation de l'entreprise tout en s'affranchissant de certaines contraintes techniques, notamment celles liées à la localisation géographique des équipements.

En fait, les VLAN introduisent la notion de segmentation virtuelle en constituant des sous-réseaux logiques selon des critères prédéfinis (ports, adresses MAC, adresses réseau...). Un logiciel d'administration permet

d'affecter chaque système raccordé à un commutateur à un réseau logique d'appartenance.

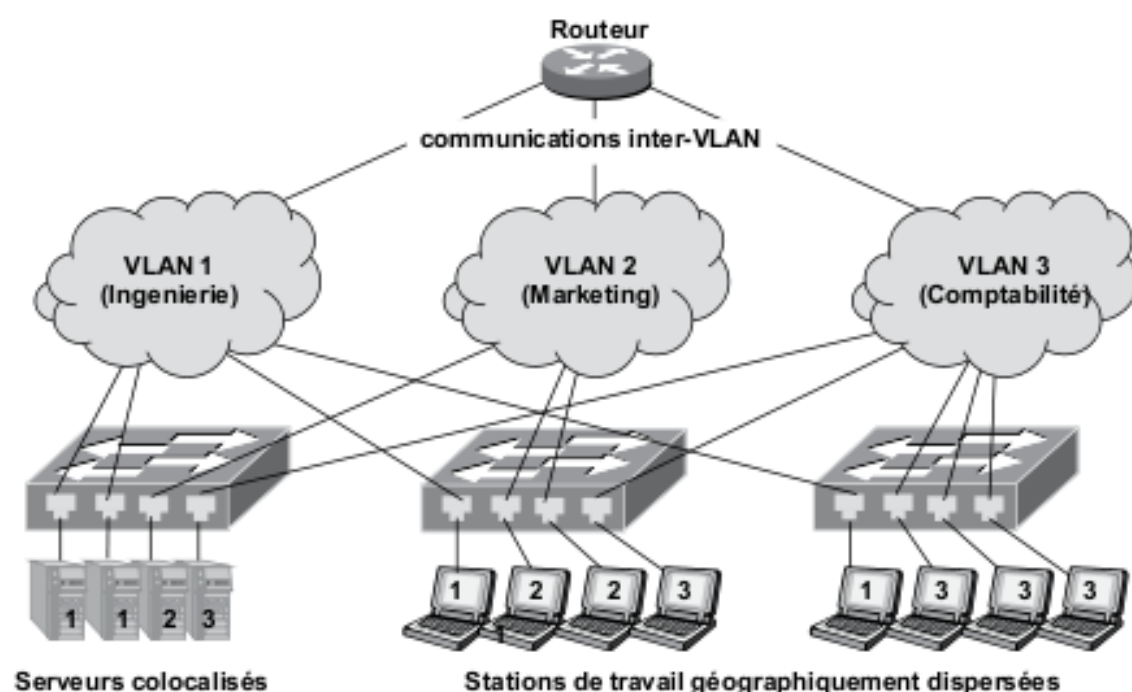


Figure 21.3 La communication inter-VLAN.

21.4 L'identification des VLAN (802.1Q)

21.4.1 Principe

Lorsqu'un réseau comporte plusieurs commutateurs, chaque commutateur doit pouvoir localiser toutes les machines (table d'acheminement) et connaître le VLAN d'appartenance de la source et du destinataire (filtrage de trafic). Lorsque le réseau est important, les tables peuvent devenir très grandes et pénaliser les performances. Il est plus efficace d'étiqueter les trames. L'étiquette identifie le VLAN de la station source, le commutateur n'a plus alors qu'à connaître les VLAN d'appartenance des stations qui lui sont raccordées.

21.4.2 La norme IEEE 802.1p/Q

Un VLAN correspond à un domaine de *broadcast*. Cependant, lorsque plusieurs VLAN sont définis sur un même segment, cette définition est mise en défaut. Il est évidemment possible d'imaginer que le commutateur transforme le *broadcast* en une rafale d'*unicast*. La solution adoptée par l'IEEE est toute différente : un seul VLAN peut être déclaré par port¹, sauf pour les liaisons inter-commutateur supportant le trafic de VLAN différents (liens dits : *trunk link*, figure 21.4).

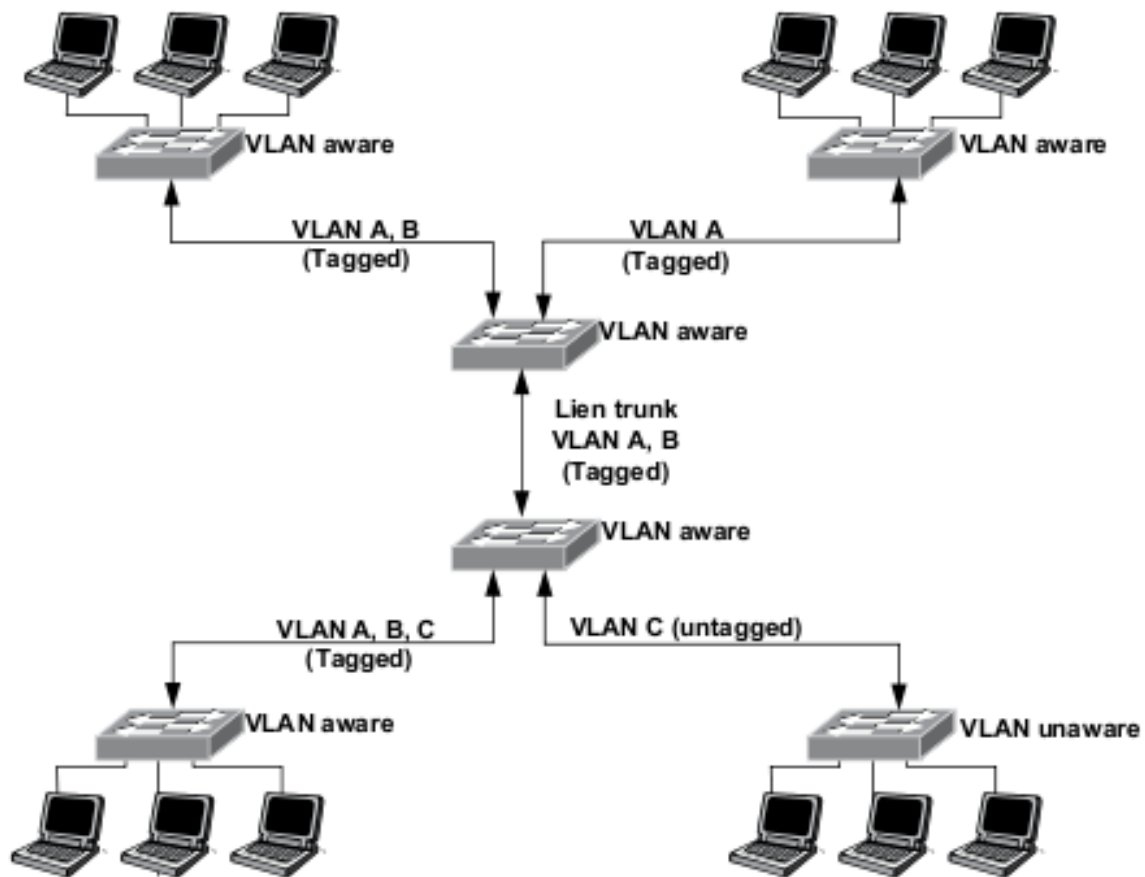


Figure 21.4 Principe de l'étiquetage des trames dans les VLAN.

¹ Certaines implémentations autorisent le raccordement de périphériques partagés par plusieurs VLAN (superposition de ports).

Les VLAN sont définis dans les normes 802.1Q et 802.1p (802.1p/Q¹) qui introduisent quatre octets supplémentaires dans la trame MAC. Ces quatre octets permettent d'identifier les VLAN (*VLAN tagging*) et de gérer huit niveaux de priorité (*Qualité of Service, QoS*). La figure 21.5 illustre l'étiquetage d'une trame MAC des réseaux de type 802.3.

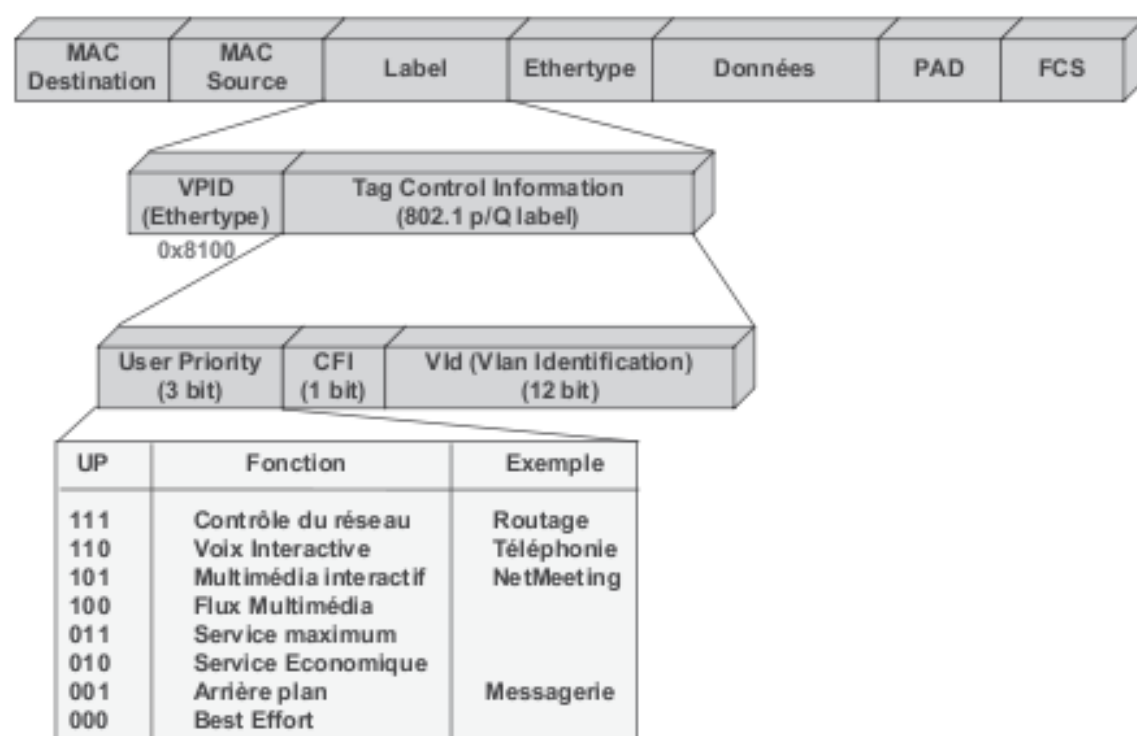


Figure 21.5 Format de la trame 802.1p/Q.

1 802.1Q concerne les VLAN, 802.1p la qualité de service.

22

L'Ethernet sans fil

22.1 Généralités

S'affranchissant d'une infrastructure câblée et autorisant la mobilité, les réseaux sans fil, sous des appellations génériques différentes, sont en plein essor. On distingue :

- ▶ les **WPAN** (*Wireless Personal Network*), de la simple liaison infrarouge à 100 kbit/s au Bluetooth initialement à 1 Mbit/s, ces technologies peu coûteuses sont essentiellement utilisées pour raccorder un périphérique informatique (imprimante...), un agenda électronique...
- ▶ les **WLAN** (*Wireless Local Area Network*), prolongent ou remplacent un réseau local traditionnel. Ces réseaux connaissent un développement important. Ils autorisent des débits allant de 2 à 200 Mbit/s (802.11n) voire plus ;
- ▶ les **WMAN** (*Wireless Metropolitan Area Network*) utilisés pour l'accès aux réseaux d'infrastructure (boucle locale radio), ils offrent des débits de plusieurs dizaines de Mbit/s ;
- ▶ enfin, les **WWAN** (*Wireless Wide Area Network*), recouvrent essentiellement les réseaux voix avec leurs extensions données (GSM, GPRS et UMTS).

22.2 La problématique de l'accès aux réseaux sans fil

Dans les réseaux à diffusion tels que les réseaux locaux, tout le monde doit pouvoir écouter tout le monde, ce qui implique l'utilisation d'une

fréquence unique. Compte tenu que l'on ne peut émettre et recevoir en même temps sur une même fréquence (éblouissement de l'antenne de réception par le signal émis), il n'est pas possible de contrôler que le signal émis n'est pas pollué par une autre émission. La détection de collision, telle que nous la connaissons dans le CSMA/CD, n'est pas réalisable, le canal radio monofréquence est un canal *half duplex*.

Dans les réseaux sans fil, pour prévenir les collisions, avant d'émettre, une station écoute le support durant un temps supérieur au délai de propagation le plus long dans le réseau. Si le support est occupé, l'émission est différée, sinon la station poursuit son écoute pendant un intervalle de temps aléatoire et, si le support est toujours libre, émet. Cette technique dite CSMA/CA pour *Collision Avoidance* est illustrée figure 22.1.

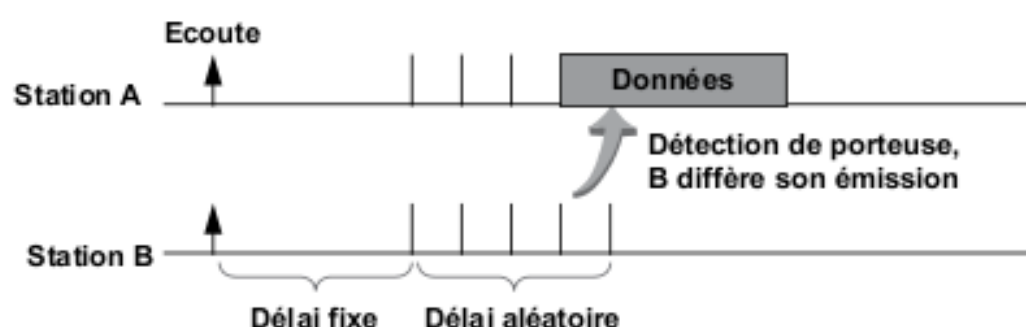


Figure 22.1 Principe du CSMA/CA.

22.3 L'architecture générale des réseaux sans fil

22.3.1 Les réseaux « ad-hoc »

Les réseaux « ad-hoc » s'affranchissent de toute infrastructure (figure 22.2). La communication est établie directement de machine à machine. Une machine peut éventuellement servir de relais pour diffuser un message vers une station non vue (au sens électromagnétique du terme) par la station d'origine (relais, routage). Les stations communiquant en mode « ad-hoc » forment un réseau appelé IBSS (*Independent Basic Service Set*).

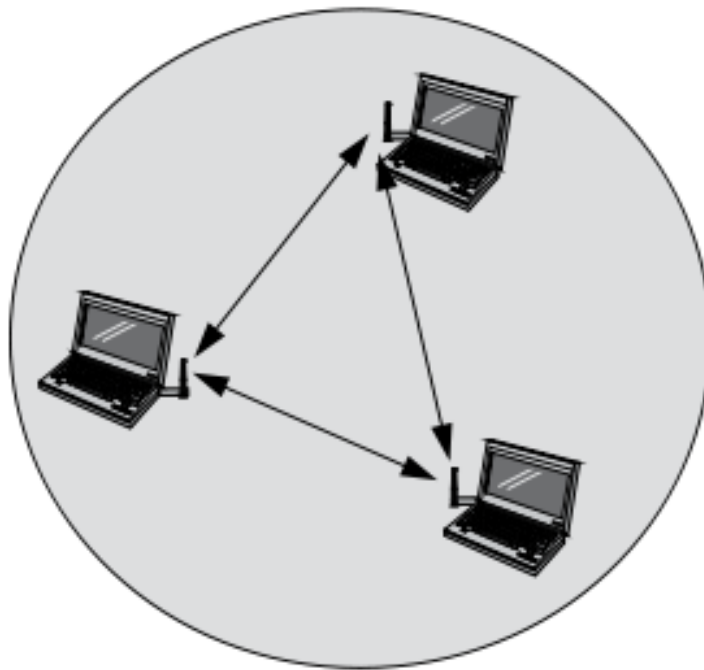


Figure 22.2 Microcellule « ad-hoc » ou IBSS.

Les réseaux 802.11 utilisent le mode infrastructure décrit ci-dessous.

22.4 Les réseaux 802.11

22.4.1 Architecture des réseaux WiFi

Le réseau IEEE 802.11 est basé sur une architecture de type cellulaire (figure 22.3). Chaque cellule, **BSS** (*Basic Service Set*), est contrôlée par une base radio, (**AP**, *Access Point*). Un BSS est identifié par un **BSSID** (*Basic Service Set Identifier*), cet identifiant, sur 6 octets, correspond à l'adresse MAC du point d'accès. Le réseau peut comporter une ou plusieurs cellules autonomes ou être le prolongement d'un réseau Ethernet traditionnel. Dans ce dernier cas, les différents points d'accès sont reliés à un réseau de distribution qui fait office de *backbone* (**DS**, *Distribution System*). La liaison entre les différents AP peut être filaire ou radio (**WDS**, *Wireless Distribution System*). L'ensemble forme un seul réseau 802.11 désigné sous le terme de **ESS** (*Extended Service Set*). L'ESS est identifié par un **ESSID**, identifiant sur 32 octets qui sert de nom au réseau. La connaissance de

l'ESSID est nécessaire pour se connecter au réseau. Les réseaux 802.3 et 802.11 sont interconnectés par un élément actif assurant l'adaptation des formats : le portail (pont à translation).

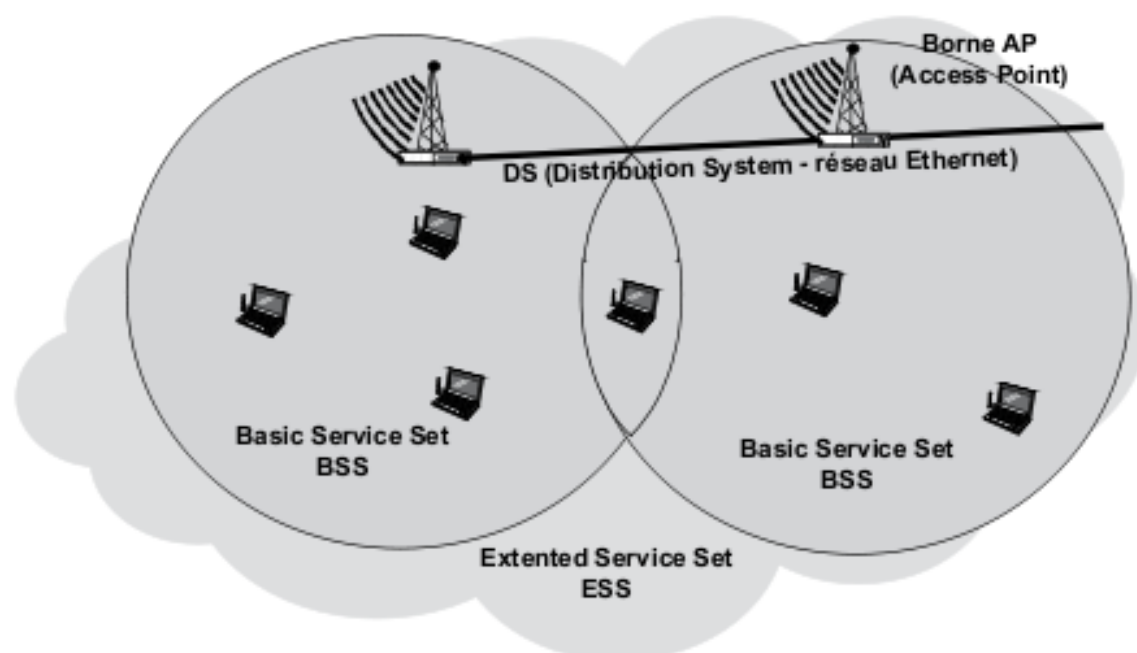


Figure 22.3 Architecture matérielle du réseau IEEE 802.11.

22.4.2 Le niveau physique ou l'interface radio

■ Généralités

Le niveau physique découpé en deux couches, l'une de convergence (**PLCP**, *Physical Layer Convergence Protocol*) et l'autre dépendant du type de modulation utilisé (**PMD**, *Physical Medium Dependent*), a pour objet la transmission des bits sur le support hertzien. Le support hertzien étant un support sujet aux perturbations électromagnétiques, aux interférences avec d'autres systèmes radios, à l'évanouissement dû aux multi-trajet (*fading*), les réseaux Wi-Fi utilisent des techniques de modulation présentant une bonne résistance à ces phénomènes perturbateurs comme les techniques d'étalement de spectre dans la bande des 2,4 GHz ou une modulation de type **OFDM** (*Orthogonal Frequency Division Multiplexing*) dans la bande des 5 GHz.

La figure 22.4 présente les différentes implémentations des réseaux 802.11, les techniques IR (infrarouge) sont aujourd'hui obsolètes.

802.11 (Wireless Local Area Networks, WLAN)				
IR (Infra-Rouge)	2,4 GHz (FHSS) (Frequency Hopping Spread Spectrum)	2,4 GHz (DSSS) (Direct Sequence Spread Spectrum)		5 GHz (OFDM) (Orthogonal Frequency Division Multiplexing)
	802.11 IR (1 ou 2 Mbit/s)	802.11 DSSS (1 ou 2 Mbit/s)	802.11b 5,5 ou 11 Mbit/s	802.11a 6, 12 ou 24 Mbit/s 9, 18, 36 ou 54 Mbit/s (options)
Autres publications:				
802.11g	Version à 2,4 GHz et 54 Mbit/s compatible 802.11b			
802.11h	Version à 5 GHz et 54 Mbit/s compatible avec 802.a			
802.11j	Version à 5 GHz et 54 Mbit/s compatible avec 802.11a et normes japonaises			
802.11n	Version à 100 voire 300 Mbit/s (2009) en rupture avec les technologies actuelles			

Figure 22.4 Synthèse des publications 802.11.

■ Les techniques d'étalement de spectre

□ L'étalement à sauts de fréquence (FHSS)

Dans la technique d'étalement de spectre à sauts de fréquence (**FHSS**, *Frequency Hopping Spread Spectrum*), la bande des 2,4 GHz est divisée en sous-canaux de 1 MHz de largeur. Cette technique se caractérise par des changements de fréquence parfaitement synchronisés (sauts de fréquence).

□ L'étalement à séquence directe

L'étalement de spectre à séquence directe (**DSSS**, *Direct Sequence Spread Spectrum*) autorise chaque canal à utiliser l'intégralité de la bande de fréquences allouée mais chaque communication utilise un code unique (**CDMA**, *Code Division Multiple Access*). Le principe est relativement simple, en multipliant la séquence binaire à émettre par une autre séquence binaire dite **Barker** telle qu'à 1 bit du signal binaire correspond N « bits » du signal émis. La nouvelle séquence ainsi constituée, dénommée *chip* (puce) ou *chipping*, présente un nombre de transitions électriques plus élevé et par conséquent un spectre plus large. En réception, le signal reçu est multiplié par la même séquence, on retrouve ainsi le signal binaire d'origine, cette opération « re-concentre » le spectre du signal d'origine et étale le spectre des éventuels signaux parasites, améliorant ainsi la résistance au bruit du signal d'origine (figure 22.5).

■ La modulation OFDM

La modulation **OFDM** (*Orthogonal Frequency Division Multiplexing*) ou encore **DMT** (*Discrete MultiTone modulation*) utilisée notamment dans les techniques **DSL** (*Digital Subscriber Line*) repose sur le principe du multiplexage fréquentiel. Le canal de transmission est découpé en sous-canaux, chaque sous-porteuse transporte N bits ou symboles. L'OFDM en répartissant le flux binaire sur M porteuses (figure 22.6), divise par M la rapidité de modulation de chaque porteuse réduisant ainsi les effets de l'interférence de symboles et optimisant l'utilisation du spectre radio-fréquence.

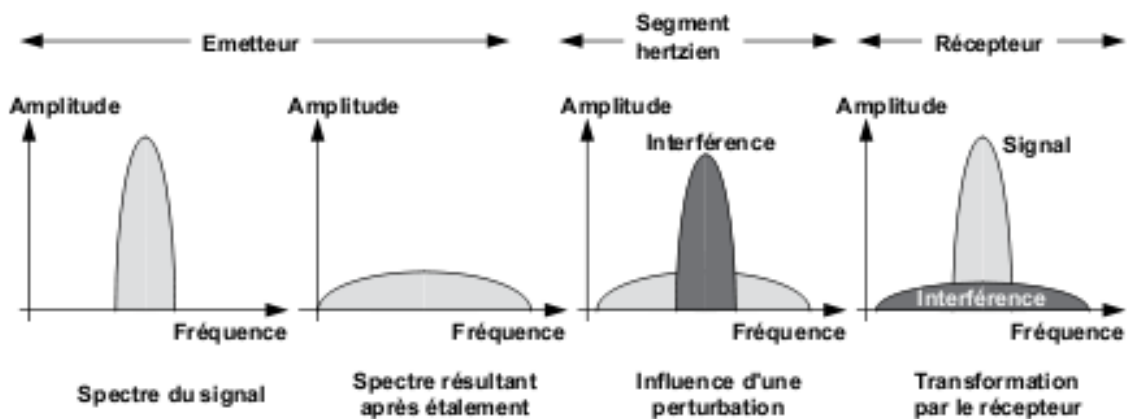


Figure 22.5 Principe de l'étalement de spectre.

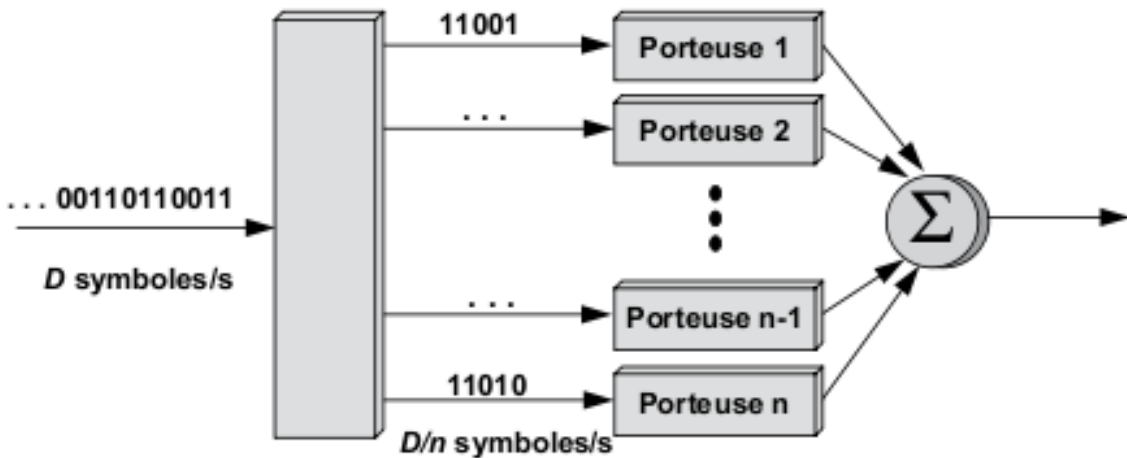


Figure 22.6 Principe de la modulation OFDM.

L'optimisation de l'occupation de l'espace fréquentiel conduit à définir des porteuses proches les une des autres, engendrant ainsi un risque d'interférence entre porteuses adjacentes (**ICI**, *Inter Carrier Interférence*). Aussi, pour minimiser ce risque, les sous-porteuses sont définies de telle manière que le maximum de puissance de leur spectre corresponde au minimum de puissance des porteuses voisines (porteuses dites orthogonales).

■ La technologie MIMO

Les ondes électromagnétiques se réfléchissent sur les obstacles engendrant plusieurs canaux de communication et générant des échos parasites. Les réseaux Wi-Fi, destinés à être utilisés en milieu fermé, n'échappent pas à cette règle, aussi un récepteur reçoit-il de nombreuses fois le même signal

(diversité spatiale), la technologie MIMO (*Multiplexing In Multiplexing Out*) met à profit cette technique de diversité spatiale (SDM, *Spatial Diversity Multiplexing*) pour exploiter non seulement les ondes directes mais aussi les échos améliorant ainsi le débit possible et la portée du système.

22.4.3 Le niveau MAC ou l'accès au support

■ Généralités

La couche MAC définit deux modes d'accès au canal, un mode d'accès par défaut qui organise un accès à compétition (**DCF**, *Distribution Coordination Function*) et un mode d'accès optionnel proche du mode à réservation (**PCF**, *Point Coordination Function*) avec un contrôle centralisé géré par le point d'accès. Ces modes d'accès mettent en œuvre un mécanisme d'accès au support contrôlé par l'utilisation de silences inter-trames (**IFS**, *InterFrame Spacing*) décrit figure 22.7. Cette technique introduit un mécanisme de priorité à l'émission d'une trame.

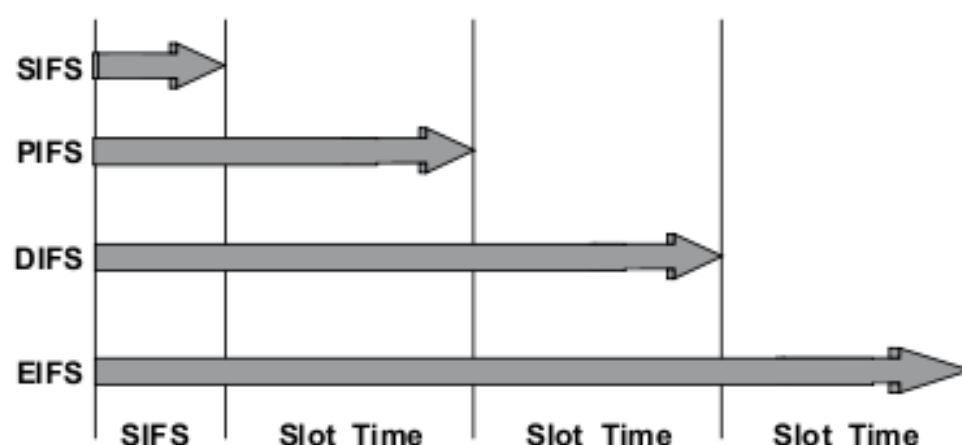


Figure 22.7 Les différents intervalles d'espacement entre trames (IFS).

Le *Short IFS* (**SIFS**) est le plus petit silence inter-message. Il donne une priorité absolue à certains messages et notamment aux acquittements. Le *Point Coordination Function IFS* (**PIFS**) est utilisé par les bornes d'accès pour l'émission de données, ce qui leur octroie un accès prioritaire par rapport à l'émission des stations toujours précédée, sauf pour les ACK, d'un silence inter-message plus important (*Distributed Coordination Function*, **DIFS**). Enfin, **EIFS** (*Extended IFS*) est utilisé pour toute retransmission

après réception d'une trame incorrecte (CRC ou collision). Chaque IFS correspond à la valeur de son précédent augmentée du *Slot_Time* tel que défini dans l'algorithme de *Backoff*.

■ Le mode Distributed Coordination Fonction (DCF)

□ Mode de base

Le mode DCF implémente une fonction de contrôle de l'accès au support distribué dont le principe de base est très simple : une station voulant transmettre un message s'assure durant un certain temps fixe et prédéterminé (DIFS) que le support est libre, puis poursuit cette écoute durant un temps aléatoire (fonction de *BackOff*) de $N \text{ Slot_Time}$, si le support est toujours libre, la station émet. Chaque message transmis doit être acquitté immédiatement, le destinataire, procède de même pour l'envoi de l'ACK, mais ne vérifie la disponibilité du support que durant un intervalle de temps réduit (SIFS). Le dialogue entre la station A et B de la figure 22.8 illustre le mécanisme.

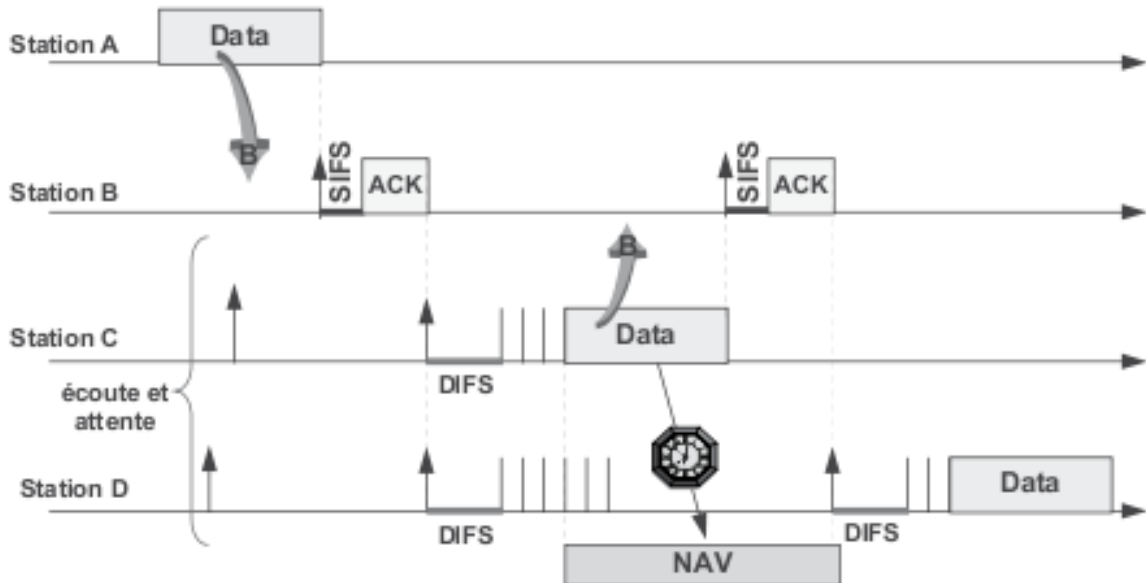


Figure 22.8 Principe de l'accès DCF, sans collision.

Pour réduire la probabilité de collisions, ce mécanisme est complété par un mécanisme de réservation dit : *Virtual Carrier Sense*. Toute trame émise contient une information sur la durée totale du cycle de

transmission (données + ACK). Ainsi, toute station écoutant le support reçoit cette information, elle positionne alors un temporisateur (NAV, *Network Allocation Vector*) et s'interdit toute émission durant cet intervalle de temps : dans la figure 22.8, le dialogue entre C et B est entendu par D qui positionne son NAV et diffère son émission d'autant. À l'échéance de ce timer, D réinitialise un cycle d'acquisition du support (DIFS + *BackOff*).

Ce mécanisme tente de prévenir les collisions, mais ne garantit nullement que deux stations n'émettent pas en même temps. La collision ne pouvant être détectée par écoute du support, l'émission des données se poursuit, c'est la non-réception de l'ACK qui informe la station de l'état de collision. La station doit alors retransmettre le message dans son intégralité. Cependant, durant l'état de collision, les stations n'ont pu positionner leur NAV (message altéré), aussi, pour éviter des collisions multiples, le délai d'écoute du support est-il étendu (EIFS).

■ Le mode Point Coordination Function (PCF)

Le mode DCF introduit un mode d'accès à compétition, l'acquisition du support n'est pas bornée et, par conséquent, ce mode de transmission ne convient pas aux données ayant des exigences temporelles strictes comme les flux multimédias (données isochrones). Aussi, la norme introduit-elle un mode d'accès optionnel : le PCF, *Point Coordination Function*. Ce mode définit un mécanisme de scrutation (*Polling*) géré par le point d'accès (AP). Cependant, le mécanisme de *polling*, s'il borne l'accès, est consommateur de bande passante.

■ Format de la trame 802.11

La trame 802.11 comporte trois parties, un en-tête MAC (*MAC header*) de 30 octets, un champ Données variant de 0 à 2 312 octets et enfin, le traditionnel champ de contrôle : le FCS (*Frame Check Sequence*). L'émission de la trame est toujours précédée par celle d'un préambule de synchronisation de 12 octets (format court) ou 16 octets (format long), suivi d'un en-tête ajouté par la couche de convergence (PLCP, *Physical Layer Convergence Protocol*). La figure 22.9 détaille chacun de ces champs.

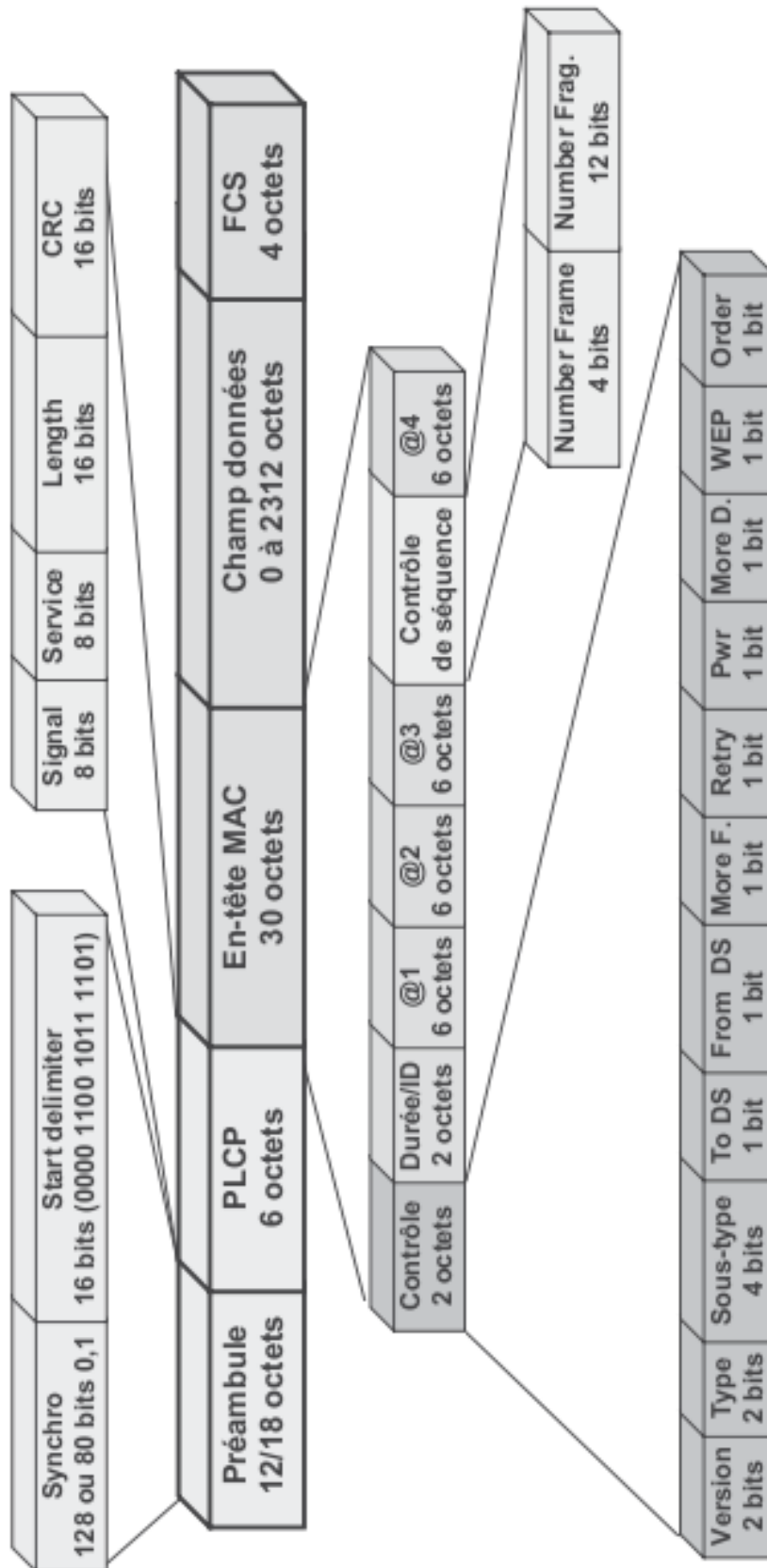


Figure 22.9 Trame MAC 802.11.

L'en-tête physique de convergence contient un ensemble d'informations nécessaire au décodage de la trame et en particulier le champ **Signal** qui indique le débit d'émission de la trame et le champ **Length** qui fournit en octets la longueur de la trame. L'ensemble étant protégé par un CRC. Le champ **Service** est inutilisé (positionné à 0). Le champ **Contrôle** sur 2 octets fournit les informations suivantes :

- ▶ **Version**, sur 2 bits, ce champ est actuellement positionné à 0 ;
- ▶ **Type et Sous-Type** distinguant les différentes trames. ;
- ▶ **To Ds et From Ds**, ces champs utilisés uniquement en mode infrastructure (*Distribution system, Ds*), permettent d'indiquer (bit à 1) si la trame est destinée au point d'accès (*To Ds*) ou s'elle provient du point d'accès (*From Ds*) ;
- ▶ **More Fragment**, ce champ indique si 1 fragment suit celui en cours (bit à 1) ou si la trame reçue est le dernier fragment (bit à 0) ;
- ▶ **Retry**, ce champ à 1 indique que la transmission en cours est en fait une retransmission sur non réception d'ACK de niveau MAC ;
- ▶ **Pwr (Power)**, ce champ demande à la station de passer en mode économie d'énergie ;
- ▶ **More Data**, en mode économie d'énergie, ce bit est positionné par le point d'accès pour indiquer qu'il dispose encore de données à transmettre après la trame en cours ;
- ▶ **WEP**, ce bit indique que le champ Données de la trame est chiffré avec l'algorithme du WEP ;
- ▶ Enfin, le champ **Order** indique que la trame reçue appartient à une classe de service où l'ordre des données doit être respecté (*Strictly Ordered Class Service*).

Le champ **Durée/ID** a deux significations. En mode économie d'énergie, dans les trames de *polling*, il identifie la station « pollée ». Dans toutes les autres trames, il indique la durée d'utilisation du canal, et permet le positionnement du NAV (*Network Allocation Vector*).

Le sous-champ **Contrôle de séquence** assure la numérotation des trames émises (*Number frames*). Le second sous-champ (*Number fragment*) est utilisé par le destinataire pour assurer le réassemblage des différents fragments.

■ Notion de sécurité dans les réseaux Wi-Fi

Les ondes électromagnétiques se propagent indépendamment de tout support, elles peuvent être reçues par toute station à l'écoute, aussi se prémunir contre les écoutes clandestines est l'une des préoccupations majeures de tout système de transmission sans fil (faisceaux hertziens, téléphonie mobile...). La norme 802.11 prévoit la possibilité de chiffrer les données (**WEP**, *Wired Equivalent Privacy*), mais les nombreuses faiblesses du chiffrement RC4 ont conduit le comité 802.11 à éditer un nouveau standard (802.11i, **WPA**, *Wifi Protected Access*).

22.5 Conclusion

Le succès d'Ethernet réside essentiellement dans son adaptabilité aux besoins. Les différentes versions ont su offrir aux utilisateurs les débits requis par les nouvelles applications. Les VLAN ont apporté un service de sécurité que les versions 1GE et 10GE ont su prolonger à l'extérieur du LAN¹. Avec un débit maximal de 10 Gbit/s aujourd'hui et 40 voire 160 demain, Ethernet semble être la solution de transport de niveau 2 offrant ainsi un service de bout en bout sans rupture de technologie (figure 22.10). Associés à MPLS, les réseaux de demain seront vraisemblablement Ethernet de bout en bout.

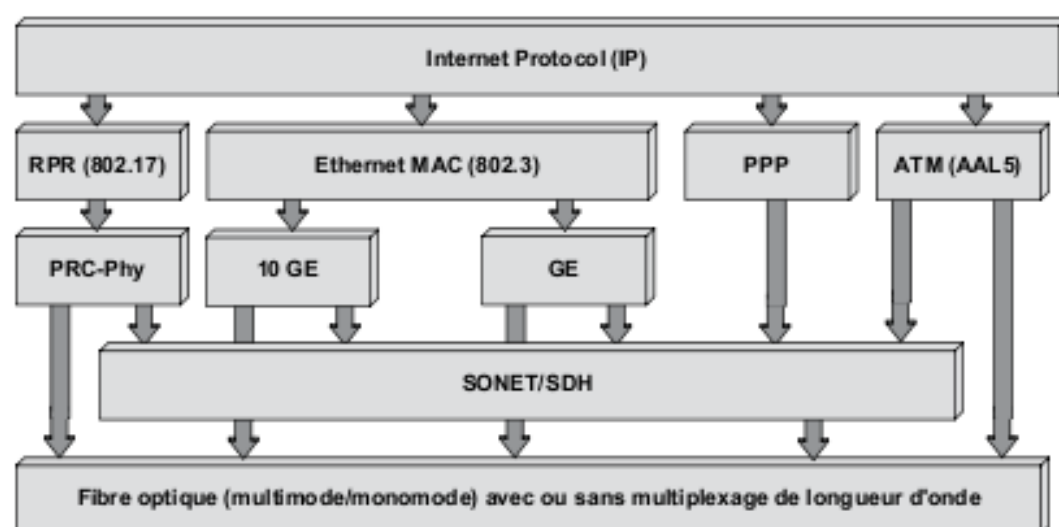


Figure 22.10 Les différentes solutions de tramage de niveau 2.

1 Ethernet Carrier Grade, voir chapitre 8, paragraphe 5.

8

Les réseaux d'opérateur

23

Structure et protocoles

Pour les entreprises, leur système d'information constitue l'un des éléments importants de leur compétitivité. Aussi, recherchent-elles des solutions fiables de mise en relation de leurs divers systèmes. Les réseaux privés d'entreprise assurent une maîtrise de la technologie mais sont d'un coût de gestion et d'évolution élevés. Aussi, les entreprises ont-elles de plus en plus recours à des opérateurs qui leur offrent connectivité, sécurité et évolutivité.

23.1 Architecture générale

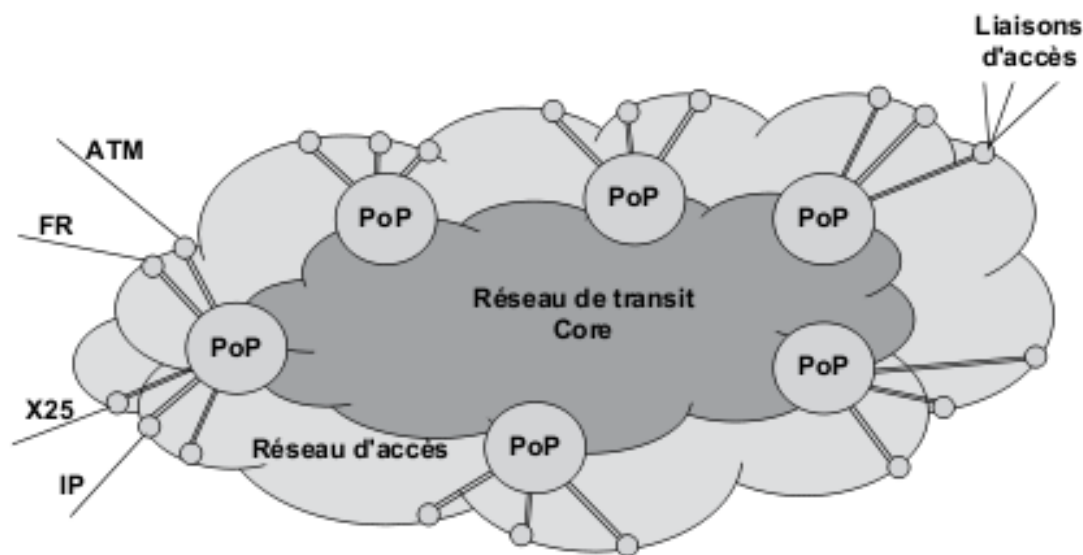


Figure 23.1 L'architecture générale des réseaux d'opérateur.

Les réseaux d'opérateur assurent deux fonctions essentielles, la collecte des flux des différentes sources par un ensemble de liens formant le **réseau d'accès** et l'acheminement de ce trafic par leur **réseau de transit**. Certains

opérateurs n'assurent que l'une des deux fonctions, on distingue alors les opérateurs de boucle locale et les opérateurs de transit. On appelle point de présence (**PoP**, *Point of Presence*) l'interface d'interconnexion entre le réseau d'accès et le réseau de transit (figure 23.1).

23.2 Structure générale d'un réseau

Un réseau peut être vu comme étant la superposition de trois plans (figure 23.2) :

- ▶ Le **plan usager**, aussi appelé plan applications, qui correspond à l'installation privée de l'utilisateur final ;
- ▶ Le **plan service** qui correspond au point où le service requis par l'utilisateur, service données ou voix, est offert par l'opérateur. Ces réseaux peuvent être privés ou publics. L'utilisateur est relié au plan service par une liaison d'abonné appelée aussi **boucle locale** ;
- ▶ Enfin, le **plan transmission** qui correspond au réseau réel de transport des flux numériques. Ce sont les techniques de numérisation et de multiplexage qui ont autorisé le transport de manière banalisée de tout type de flux : voix, données, images sur un même support.

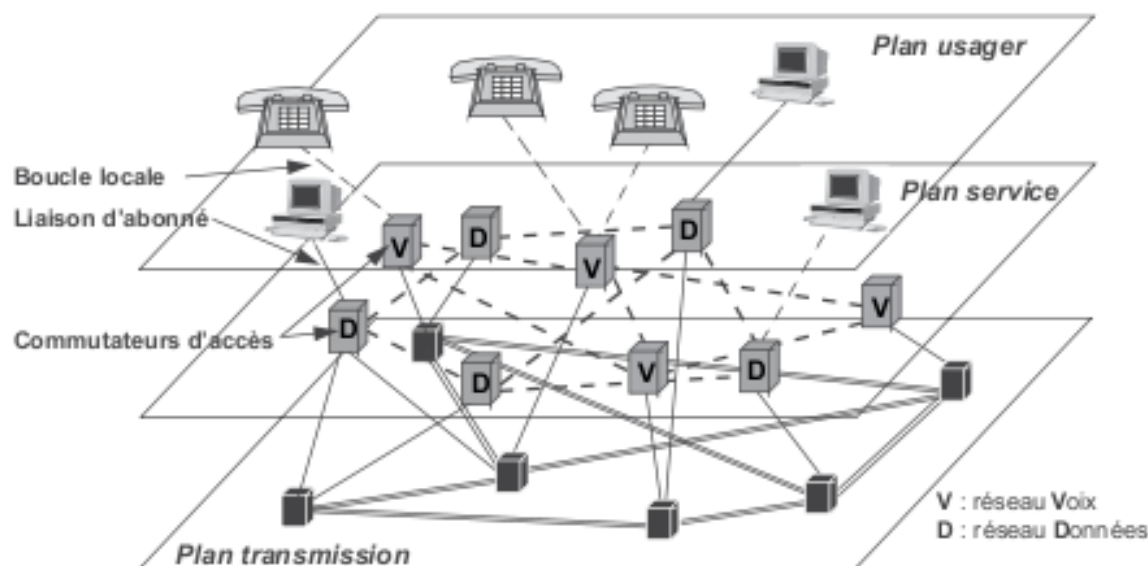


Figure 23.2 Les trois plans d'un réseau de transmission.

23.3 Le plan de transmission

Jusqu'aux années 1960, les réseaux voix et données étaient physiquement distincts, l'un analogique, l'autre numérique. La numérisation de la voix a autorisé le multiplexage des différents flux voix et données optimisant ainsi l'utilisation des supports de transmission. Deux modes de multiplexage ont été définis dans le temps, ils se distinguent essentiellement par le mode synchronisation des différents nœuds du réseau.

Dans la première technique les horloges de chacun des nœuds sont indépendantes, un nœud déduit l'horloge de son nœud amont et resynchronise le flux, c'est la hiérarchie dite **PDH** (*Plesiochronous Digital Hierarchy*) celle-ci a constitué la base de tous les réseaux de transmission jusqu'aux années 1990 et guide encore le principe de certains raccordements aux réseaux des opérateurs.

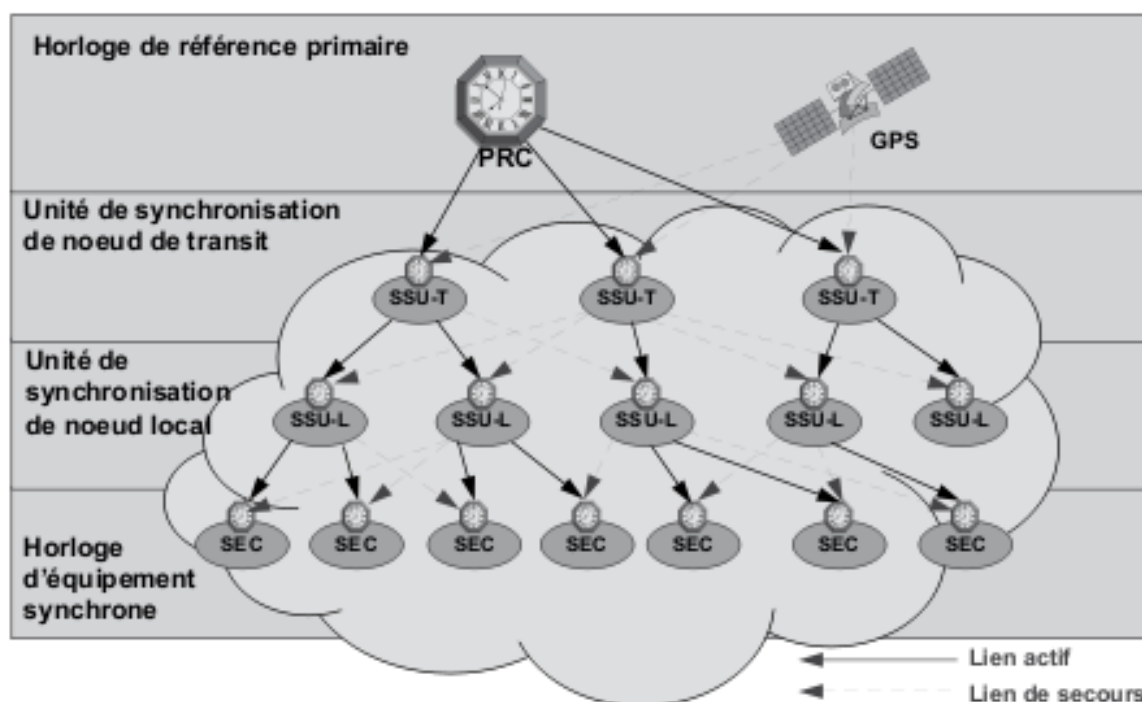


Figure 23.3 La synchronisation maître-esclave dans un réseau SDH.

Outre la rationalisation de l'utilisation des supports, la hiérarchie PDH a eu le mérite de résoudre les difficultés de synchronisation de

flux provenant de sources différentes aux horloges indépendantes mais proches (plésio). Fondée sur un réseau de distribution d'horloge, la hiérarchie synchrone (**SDH**, *Synchronous Digital Hierarchy*) garantit la délivrance des bits en synchronisme d'une horloge de référence (figure 23.3). Elle autorise des débits plus élevés, apporte des solutions d'administration et de contrôle et enfin répond à un besoin de normalisation des interfaces optiques.

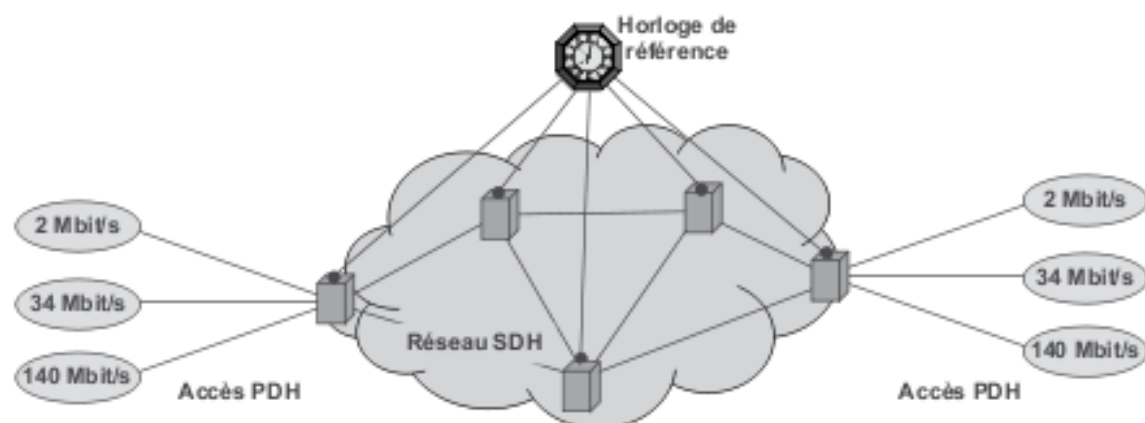


Figure 23.4 La cohabitation des techniques PDH et SDH.

Si les cœurs des réseaux sont aujourd'hui SDH, pour des raisons historiques, la distribution des débits chez l'utilisateur, hors raccordement Ethernet, repose toujours sur la hiérarchie plésiochrone (figure 23.4).

23.4 Le plan de service

Le plan de service, vu de l'utilisateur, correspond au réseau de transport de ses données. La figure 23.5 représente le réseau de transport tel qu'il est perçu par l'utilisateur.

Au cours des années 1970, la recherche de la performance a orienté les concepteurs de réseaux vers la réalisation de réseaux à commutation de paquets (*Packet switching*) en mode orienté connexion (**CONS**, *Connection Oriented Network Service*) avec le protocole X.25. Le besoin croissant de bande passante a rapidement conduit les opérateurs à rechercher des protocoles internes plus efficaces : d'abord le relais de trames (**FR**, *Frame Relay*), puis l'**ATM** (*Asynchronous Transfer Mode*).

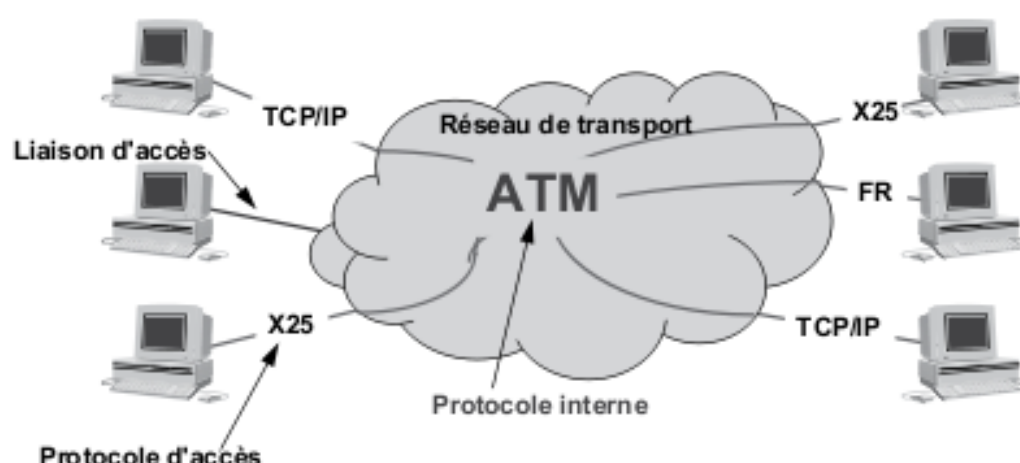


Figure 23.5 Les protocoles d'accès et protocole interne.

23.4.1 Le protocole X.25

Conçu par les PTT français, britanniques, TCTS (*Trans Canada Telephon System*) et Telnet (États-Unis), le protocole X.25, aujourd'hui considéré comme obsolète, a été le premier protocole utilisé dans les réseaux publics de données.

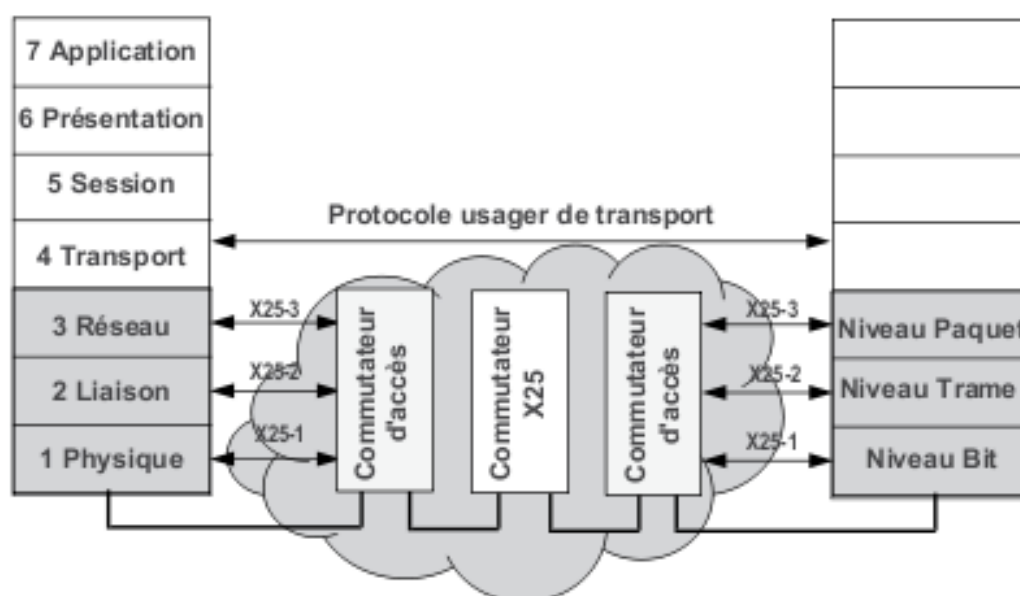


Figure 23.6 L'architecture du protocole X.25.

Le protocole X.25 couvre les trois premières couches du modèle OSI (figure 23.6) :

- La **couche physique**, niveau bit ou X.25-1 définit l'interface ETDD/ETCD. Elle est conforme à l'avis X.21 et X.21 bis de l'UIT-T ;
- La **couche liaison**, niveau trame ou X.25-2, met en œuvre un sous-ensemble d'HDLC appelé LAP-B (*High Level Data Link Control, Link Access Protocol Balanced*) ;
- La **couche réseau**, niveau paquet ou X.25-3, gère les circuits virtuels (permanents ou commutés).

■ Format des unités de données

Les paquets X.25-3 comportent les informations relatives à l'adressage : le numéro de voie logique (adressage de convention), des informations de contrôle et éventuellement des données. Les paquets sont transmis à la couche trame qui les encapsule (figure 23.7).

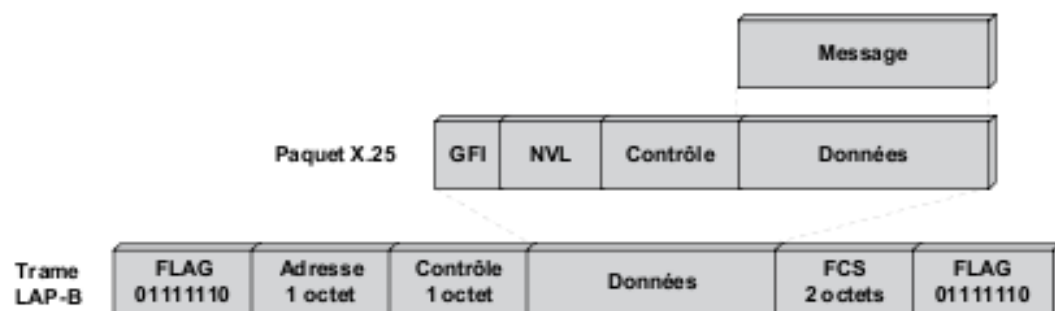


Figure 23.7 L'encapsulation des paquets X.25 dans les trames LAP-B.

Un paquet X.25, illustré figure 23.7, comporte au moins 3 octets. Le premier champ de 4 bits, dit champ **GFI** (*General Format Identifier*) définit certains paramètres de l'échange.

Le champ suivant identifie la voie logique (étiquette) : le **NVL** ou numéro de voie logique (**LCN**, *Logical Channel Number*) sur 12 bits (4 096 voies logiques identifiables) n'a qu'une signification locale.

Enfin, le dernier champ de l'en-tête est similaire au champ Commande d'HDLC. Les champs $P_{(r)}$ et $P_{(s)}$ ($N_{(r)}$, $N_{(s)}$ d'HDLC) permettent de contrôler le séquençement des paquets. Ces compteurs ne font pas double usage avec ceux du niveau trame. En effet, une connexion de niveau 2 peut multiplexer plusieurs connexions de niveau 3.

Dans le protocole X.25, la signalisation est du type dans la bande, ce qui signifie que les données et les informations de signalisation sont transportées par des unités de données similaires. Indépendamment de la pauvreté d'un tel système de signalisation, l'établissement de la route et la commutation sont réalisées par une même instance de programme. Cette architecture en multipliant les opérations d'encapsulation et de décapsulation pénalise les performances.

23.4.2 D'X.25 au Frame Relay et l'ATM

L'évolution des protocoles réseaux s'est réalisée selon deux approches :

- ▶ Le **relais de trames** ou *Frame Relay* qui correspond à un allègement du protocole HDLC version LAP-D. Ce protocole répond aux besoins de haut débit ; ne traitant pas les flux isochrones à l'origine, il a été perçu comme un protocole de transition entre X.25 et ATM ;
- ▶ Le **relais de cellules** ou *Cell Relay*, plus connu sous le nom d'**ATM** (*Asynchronous Transfer Mode*) qui utilise une technique de commutation rapide de cellules de taille fixe. ATM met en œuvre des mécanismes spécifiques pour assurer les transferts isochrones (émulation de circuits pour la voix et la vidéo).

L'augmentation du débit réel des protocoles ne peut résulter que de l'allègement des traitements intermédiaires, ce qui a conduit à :

- ▶ reporter sur les organes d'extrémité (les calculateurs) les tâches de détection et de reprise sur erreur, comme le fait TCP/IP ;
- ▶ diminuer les opérations de couches en effectuant les opérations d'acheminement (commutation) au niveau 2 ;
- ▶ formuler des hypothèses optimistes sur le comportement du réseau en n'effectuant pas de contrôle de flux entre les nœuds ;
- ▶ supprimer les acquittements intermédiaires, ceux-ci n'étant réalisés que par les organes d'extrémité (acquittement de bout en bout) ;
- ▶ simplifier le traitement dans les nœuds en n'utilisant qu'un seul type de structure de données (format) pour le transfert de données et en mettant en œuvre une signalisation par un protocole spécifique sur un

canal dédié (**canal sémaphore**) pour l'établissement des circuits et la gestion du réseau.

En particulier, l'introduction de la fibre optique dans les réseaux en fiabilisant la transmission a autorisé la simplification des protocoles (contrôle d'erreur). La figure 23.8 illustre ces concepts mis en œuvre dans le relais de trames (Frame Relay).

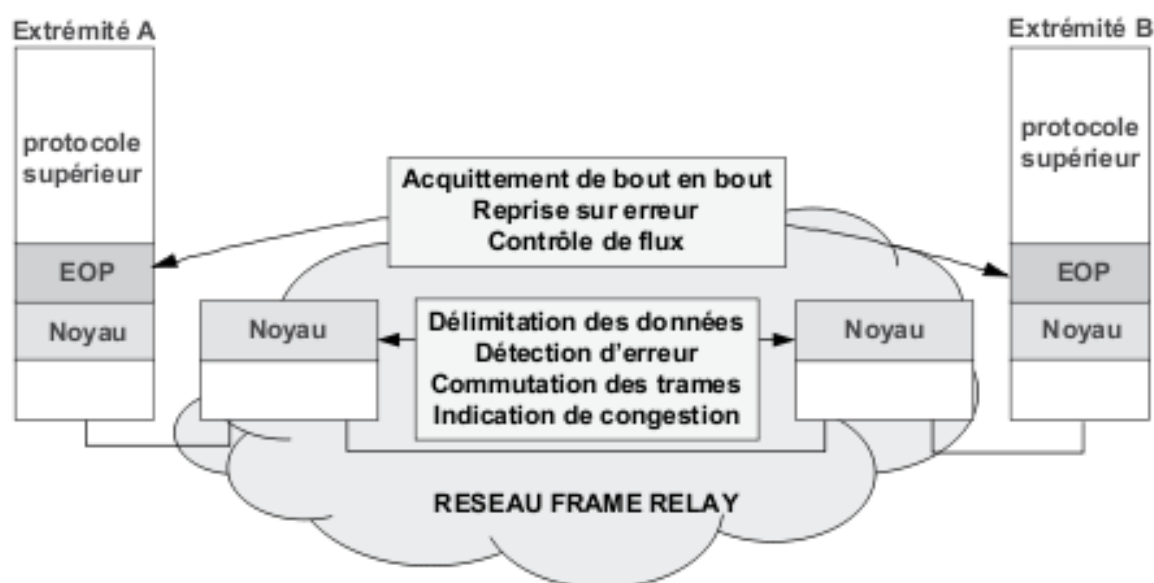


Figure 23.8 Protocoles Haut Débit et l'approche du relais de trames.

23.4.3 Introduction au Frame Relay

■ Principe de base

Le relais de trames offre un service réseau en mode connecté, la signalisation est du type canal sémaphore. Le relais de trame établit entre les équipements d'extrémité (**FRAD**, *Frame Relais Access Device*) un service de liaison virtuelle (figure 23.9). La liaison virtuelle peut être du type permanent (**PVC**, *Permanent Virtual Circuit*) ou temporaire et établi à la demande (**SVC**, *Switched Virtual Circuit*). Compte tenu de la complexité d'établissement d'un SVC et donc de la charge de calcul induite, les opérateurs n'offrent qu'un service de circuits virtuels permanents.

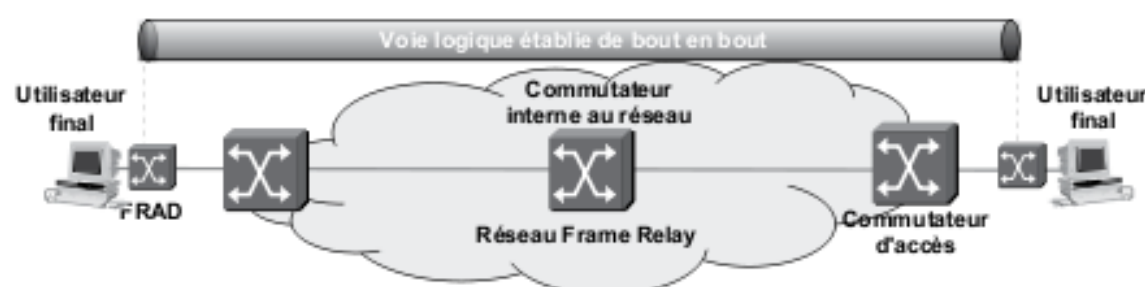


Figure 23.9 – La liaison virtuelle du relais de trames.

La liaison virtuelle est établie au niveau 3 par un protocole spécifique (signalisation hors bande), alors que les données sont acheminées au niveau 2 (commutation), ce fonctionnement est illustré par la figure 23.10.

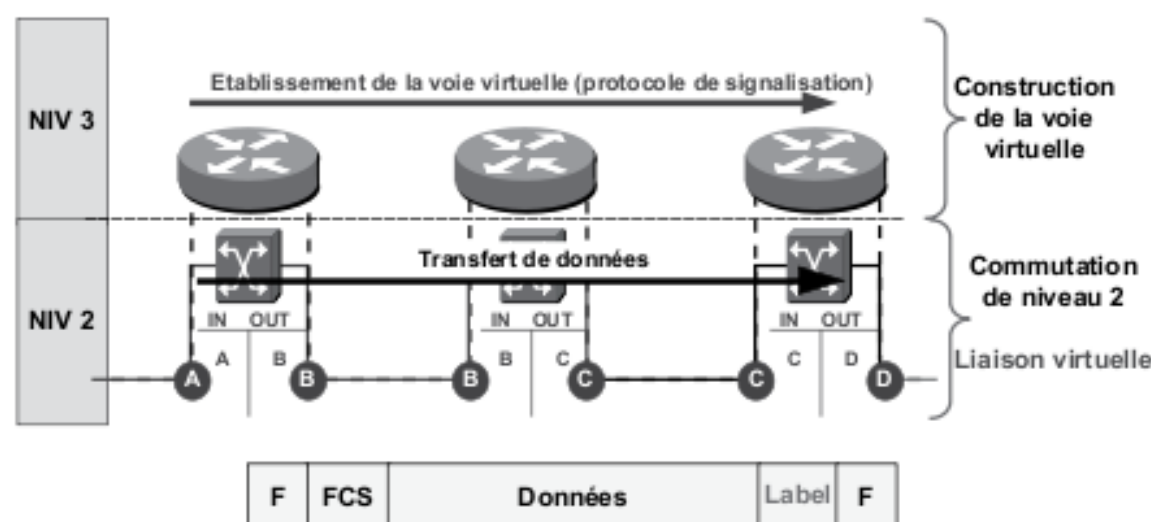


Figure 23.10 – Principe du Frame Relay.

■ Le contrôle d'accès

La simplification du protocole a conduit à la suppression de tout contrôle de flux dans le réseau et donc à fragiliser celui-ci face aux problèmes de congestion. Aussi, toute nouvelle connexion ne peut être acceptée que si le réseau est apte à la satisfaire sans préjudice pour les connexions déjà établies. Toute demande de connexion est accompagnée d'un descripteur de trafic définissant en particulier le débit moyen et le débit de pointe demandés. Un contrat de trafic est passé entre la source et le réseau (SLA, *Service Level Agreement*), il comporte un descriptif complet des paramètres de la connexion et en particulier :

- Le **CIR** (*Committed Information Rate*) ou débit moyen garanti. Le CIR caractérise le débit moyen contractuel que doit garantir le réseau. La connexion ne sera acceptée que si la somme des CIR sur le lien (ou sur le nœud) ne dépasse pas un seuil déterminé par le gestionnaire du réseau ;
- L'**EIR** (*Excess Information Rate*) ou surdébit autorisé au-dessus duquel tout bloc de données soumis au réseau est détruit ;
- Le temps d'analyse du trafic (T_c).

Le CIR définit le volume moyen admis dans le réseau ou **Bc** (*Committed Burst size*) tel que $B_c = CIR \cdot T_c$. L'EIR précise le volume maximal autorisé tel que $B_c + B_e = (CIR + EIR) \cdot T_c$ où **Be** (*Excess Burst size*) représente le volume excédentaire admis au-dessus du contrat B_c (figure 23.11).

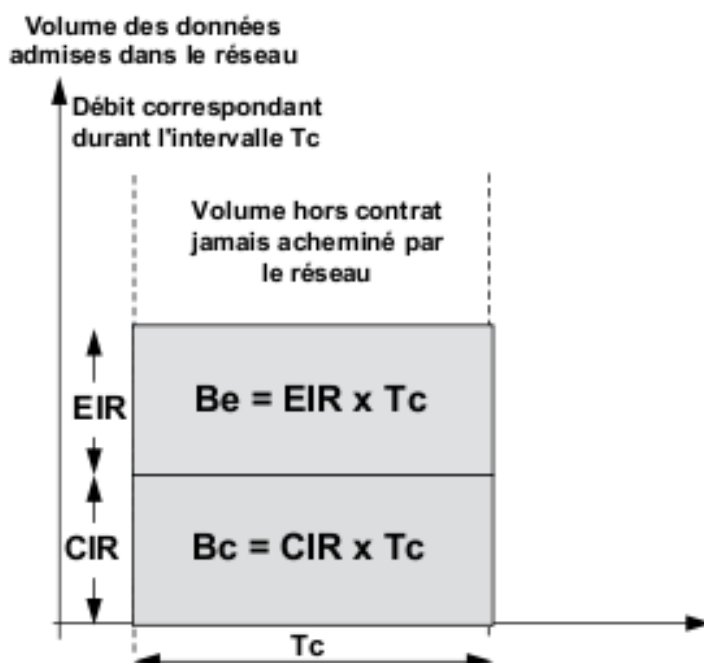


Figure 23.11 Relation entre les différents paramètres.

L'acceptation d'une nouvelle connexion ou l'établissement d'une nouvelle voie implique que chaque commutateur mémorise les différents paramètres des connexions déjà réalisées.

23.4.4 Introduction à l'ATM (Asynchronous Transfer Mode)

■ Généralités

L'ATM a été développé dans le cadre de l'évolution du réseau téléphonique (RNIS Large Bande, **B-ISDN** ou *Broadband Integrated Service Digital Network*). En effet, la commutation de circuits des réseaux voix traditionnels monopolise la bande passante alors que la commutation de paquets, en autorisant le multiplexage statistique des sources, optimise l'utilisation de la bande passante. C'est sur ces bases que le CNET (Centre national d'étude et de télécommunication) a décrit, en 1982, une technique de multiplexage asynchrone (**ATD**, *Asynchronous Time Division*) qui allait donner naissance à l'ATM.



Figure 23.12 L'origine de l'ATM.

La figure 23.12 illustre le principe d'un réseau RNIS. Le cœur de réseau est un réseau en mode paquet (cellule), le protocole ATM optimise l'utilisation des ressources. L'abonné est raccordé à ce réseau *via* un PABX (*Private Automatic Branch eXchange*), commutateur téléphonique faisant office de passerelle d'accès à ce réseau et séparant les flux voix des flux de données. La liaison est du type E1 (liaison à 2 Mbit/s composée d'un « train » de 32 IT de 8 bits 8 000 fois par seconde).

Destiné au transport de la voix, des données et de l'image, ATM est une technologie en mode connecté, les données ne sont acheminées dans le réseau qu'après l'établissement d'une voie virtuelle (**VCC**, *Virtual Channel Connection*).

L'architecture de l'ATM comporte trois couches (figure 23.13) dont les fonctions essentielles sont :

- ▶ assurer l'adaptation des cellules au système de transport physique utilisé (couche physique),
- ▶ effectuer la commutation et le multiplexage des cellules (couche ATM à proprement parler),
- ▶ adapter les unités de données (segmentation et réassemblage) des protocoles supérieurs à la couche ATM (couche **AAL**, *ATM Adaptation Layer*) et mettre en place des mécanismes spécifiques à chaque type de données transportées.

OSI	IEEE - 802	UIT - ATM
3 - Réseau		
2 - Liaison de données	LLC (802.2)	AAL Adaptation ATM Layer
	MAC	CS Convergence Sublayer SAR Segmentation And Reassemblage
		ATM Layer
1 - Physique	PMD	Physical Layer (OC3...)

Figure 23.13 L'architecture ATM.

■ La taille des unités de données ou cellules

Étudié dans le cadre du développement du RNIS large bande, ATM voit ses caractéristiques fortement conditionnées par le transfert de flux isochrone telle que la voix. Ce dernier point a été déterminant dans le choix de la taille des unités de données (cellules), celle-ci doit permettre une commutation rapide et minimiser la gigue d'insertion. La figure 23.14 montre l'effet de la taille du bloc de données, ici cellules, sur le temps de transfert dans un réseau en mode paquet (H représente la taille de l'en-tête).

Ces deux impératifs (temps de commutation et gigue) ont conduit à l'adoption d'une unité de données de petite taille. Après de nombreuses discussions, un compromis a été adopté et la taille des unités de données (cellules) a été fixée à 48 octets et 5 octets d'en-tête.

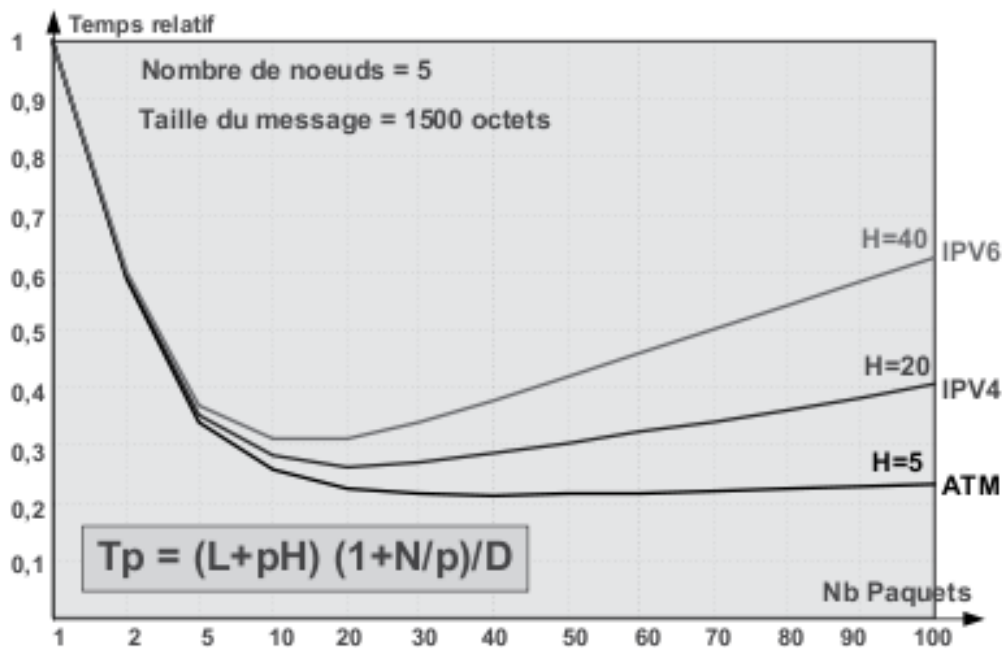


Figure 23.14 Influence de la taille des données sur le temps de transfert.

■ L'adressage dans les réseaux ATM

À l'instar de X.25 ou du Frame Relay, l'ATM utilise dans le réseau un adressage de convention identifiant les voies virtuelles (multiplexage par étiquette) et un adressage hiérarchique à la périphérie du réseau de type E.164 pour les réseaux publics (adressage du RNIS).

La taille des tables de commutation est l'un des facteurs principaux intervenant dans l'efficacité de la commutation. ATM, à des fins d'efficacité, utilise deux niveaux d'identification. En effet, dans un réseau, plusieurs sources partent d'un même commutateur d'entrée et se dirigent dans une même direction. Plutôt que de gérer les N connexions, il est plus aisé de les regrouper (identifiant secondaire) et de ne traiter dans le cœur du réseau que cet identifiant de second niveau (figure 23.15), ainsi :

- ▶ un premier niveau identifie la voie virtuelle (**VCI**, *Virtual Channel Identifier*). Il s'agit de l'identifiant des flux échangés entre deux systèmes d'extrémité, notion similaire au numéro de voie logique d'X.25 ou du DLCI du Frame Relay ;
- ▶ un second niveau regroupe (agrégation de flux), un ensemble de voies virtuelles ayant une même destination (nœud intermédiaire ou inter-

face d'utilisateur) sous un même identifiant dit *Virtual Path Identifier (VPI)*. Le VPI représente une connexion semi-permanente entre deux nœuds du réseau, elle est établie et contrôlée par l'administrateur du réseau.



Figure 23.15 La double identification d'ATM.

Les commutateurs de second niveau, appelés **brasseurs**, commutent l'ensemble des voies virtuelles affectées à un faisceau (acheminement selon les VPI). Les commutateurs sont situés en périphérie du réseau (commutateur d'accès), les brasseurs assurant la commutation des faisceaux virtuels en interne dans le réseau (figure 23.16).

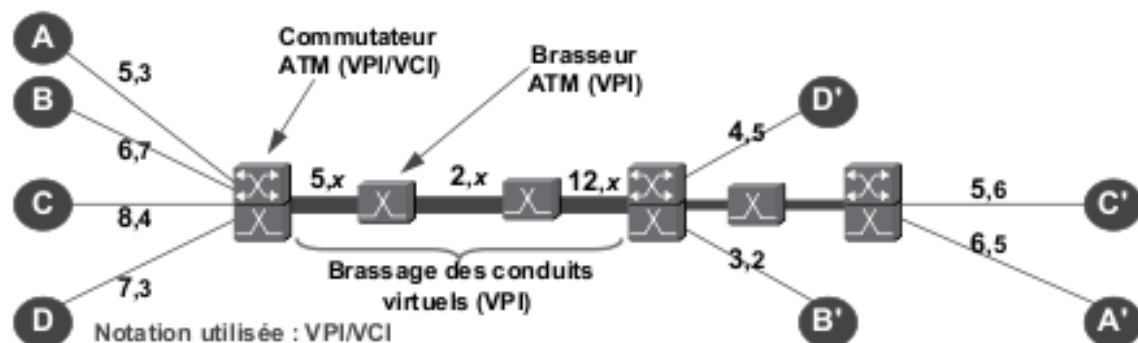


Figure 23.16 Le double niveau d'acheminement des cellules.

■ Les formats de l'en-tête de cellule ATM

À des fins d'efficacité, l'en-tête ATM ne contient que les informations strictement nécessaires à l'acheminement, au contrôle du type de données et à la protection de l'en-tête. ATM utilise deux formats d'en-tête, le premier est utilisé à l'interface usager/réseau (**UNI**, *User to Network Interface*) et le second en interne dans le réseau (**NNI**, *Network to Network Interface*). La figure 23.17 représente ces deux formats.

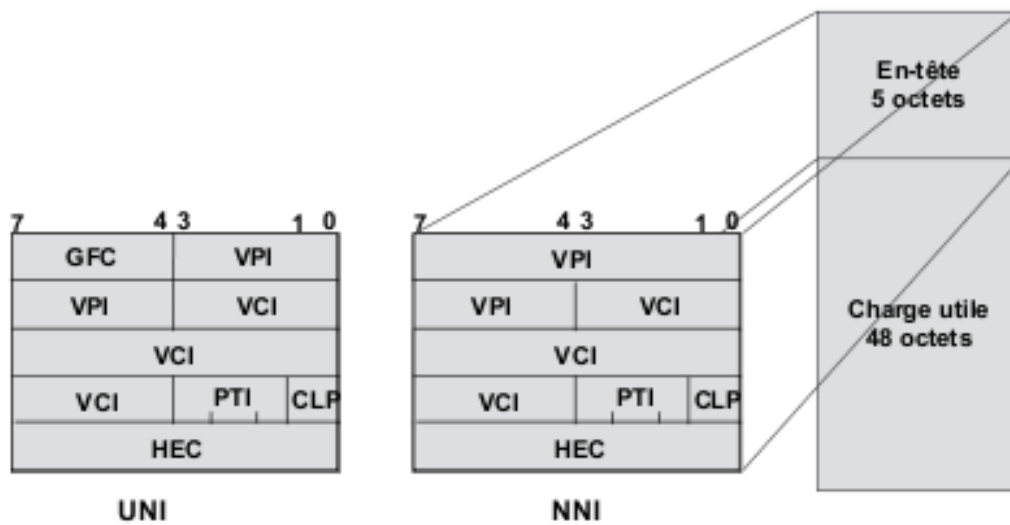


Figure 23.17 Les formats de l'en-tête de cellule ATM.

En mode contrôlé (*Controlled*), le champ **GFC** (*Generic Flow Control*) autorise le partage équitable de l'accès au réseau aux différentes stations dans une configuration point à multipoint. En mode point à point, le GFC permet la résolution des conflits d'accès (résolution des contentions) et le contrôle de flux à l'interface usager/réseau. Non utilisé, ce champ doit être mis à 0.

Les mécanismes précédents sont définis à l'interface usager, les quatre bits du GFC sont « récupérés » en cœur de réseau pour étendre le champ VPI des cellules NNI. Cette technique conduit à définir deux types d'en-tête de cellule selon que l'on se situe à l'interface usager/réseau (UNI) ou à une interface réseau/réseau (NNI) :

- ▶ les cellules UNI identifient 65 536 VCI (16 bits) et 256 VPI (8 bits) ;
- ▶ les cellules NNI identifient 65 536 VCI et 4 096 VPI (12 bits).

Le champ **PT** (*Payload Type*) sur trois bits indique le type de charge contenue dans le champ Données (Données, signalisation...).

Le bit de préférence à l'écartement (**CLP**, *Cell Loss Priority*) indique, lors d'un état de congestion, les cellules à éliminer en priorité. En fait, le bit CLP à 1 indique une cellule de priorité basse, à 0 il identifie une cellule de priorité haute.

■ Le contrôle d'erreur

Le champ HEC (*Header Error Control*) assure un contrôle d'erreur de type CRC dont la portée est limitée à l'en-tête. L'HEC ou octet de contrôle (figure 23.18) est le reste de la division booléenne (modulo 2) des quatre premiers octets de l'en-tête par le polynôme générateur : $G(x) = x^8 + x^2 + x + 1$, auquel est ajouté le polynôme $C(x) = x^6 + x^4 + x^2 + 1$ (addition booléenne soit un OU exclusif avec 01010101).

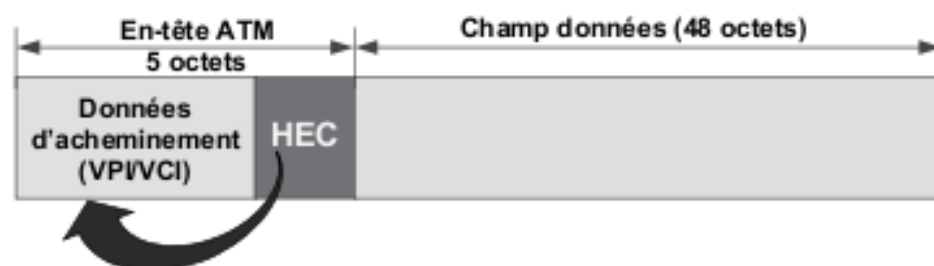


Figure 23.18 La portée du champ HEC.

En cas d'erreur, la cellule est éliminée, la reprise sur erreur ou sur temporisation est éventuellement confiée aux couches supérieures des systèmes d'extrémité (*End Systems*).

■ La délimitation des cellules

L'ATM n'utilise aucun fanion pour délimiter les cellules. Celles-ci ayant une taille fixe et une fréquence de récurrence élevée, il suffit de se positionner correctement sur un octet pour reconnaître les limites des cellules. Le HEC est calculé au fil de l'eau sur le flot de bits reçu dans un registre FIFO de 40 bits (5 octets) représenté figure 23.19. Lorsque le cinquième octet contenu dans le *buffer* correspond au résultat du calcul de l'HEC sur les 4 octets précédents, l'HEC et donc l'alignement sont supposés trouvés.

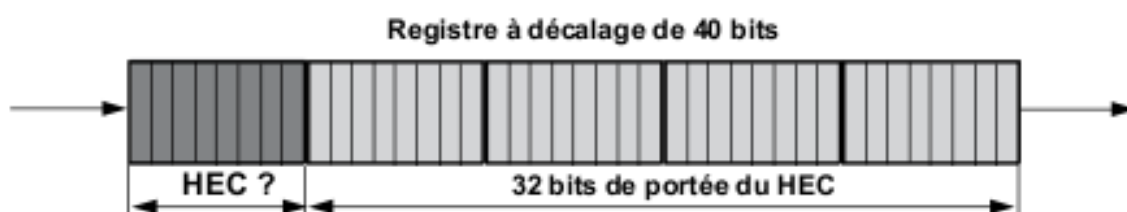


Figure 23.19 Principe du repérage du HEC.

■ Le modèle ATM

L'ATM comporte trois couches dont les fonctionnalités principales sont décrites ci-après. Les relations entre les différentes couches sont représentées en figure 23.20.

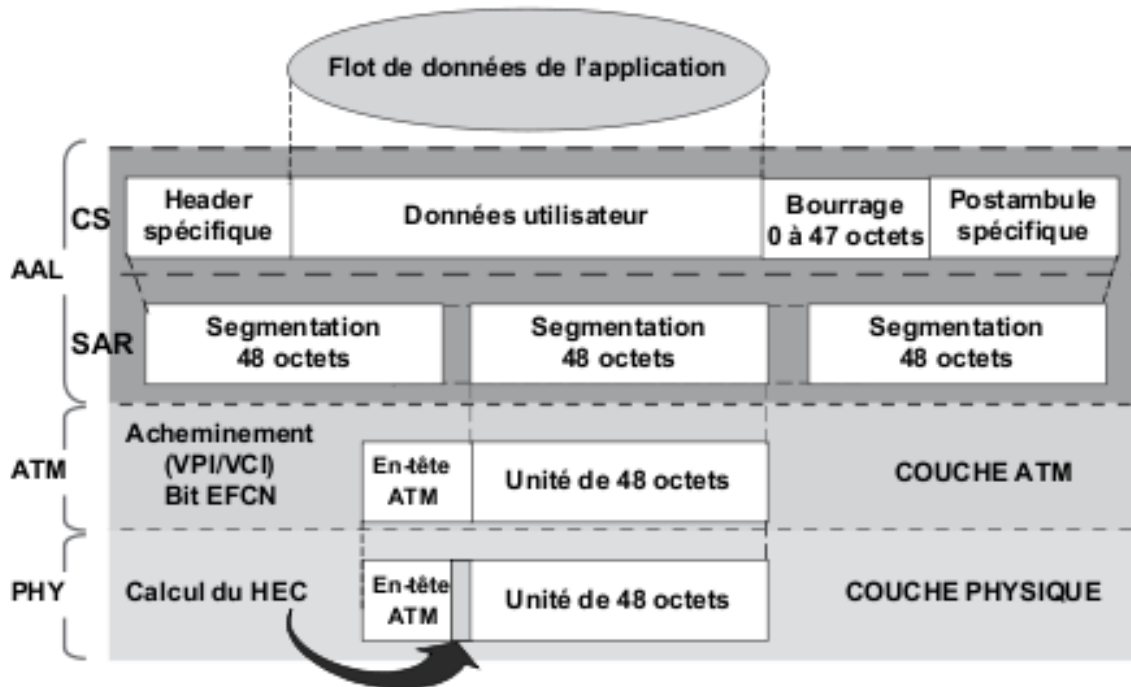


Figure 23.20 Relations entre les différentes couches de l'ATM.

La couche **ATM** est indépendante du sous-système de transport physique et des services d'applications qui l'utilisent. Elle assure les fonctions de multiplexage et démultiplexage des cellules, la génération et l'extraction des en-têtes, l'acheminement (commutation) des cellules et la translation des VPI/VCI enfin, le contrôle de flux (*GFC, Generic Flow Control*) à l'interface UNI (*User Network Interface*).

La couche physique (**PHY**) est chargée de fournir à la couche ATM un service de transport des cellules.

Enfin, la couche **AAL** (*ATM Adaptation Layer*) garantit aux applications utilisateurs la qualité de service requise par l'application. Quatre types d'AAL sont proposés aux applications : AAL1, AAL2, AAL3/4 et AAL5 (figure 23.21). La couche AAL est divisée en deux sous-couches :

- ▶ la sous-couche de convergence (CS, *Convergence Sublayer*) est destinée à incorporer les informations spécifiques au type d'AAL utilisé ;
- ▶ la sous-couche de segmentation et de réassemblage (SAR, *Segmentation and Reassemblage*) adapte le format de l'unité de données issue de la sous-couche CS au format requis par ATM (cellules de 48 octets).

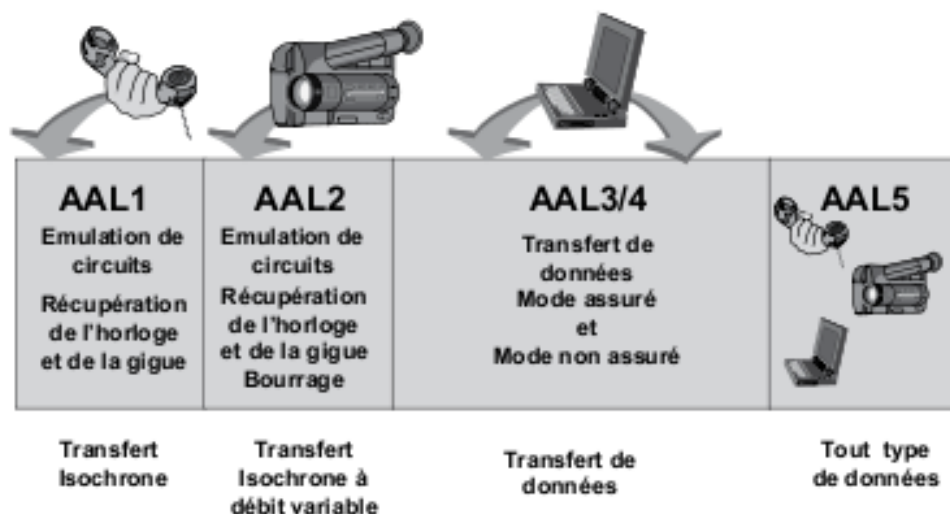


Figure 23.21 Les couches d'adaptation d'ATM.

L'AAL5, définie par l'ATM Forum pour remplacer l'AAL3/4 est aujourd'hui utilisée pour tout type de flux. La figure 23.22 représente l'encapsulation protocolaire de l'AAL5.

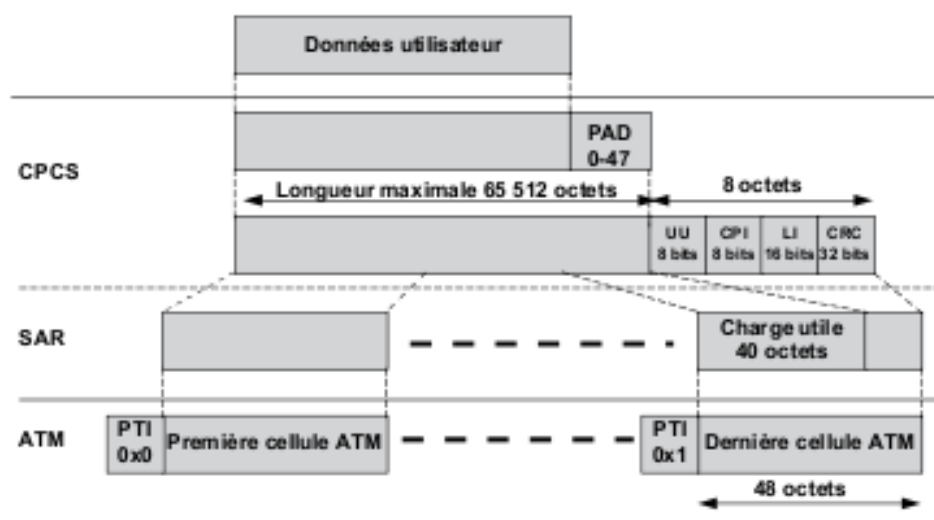


Figure 23.22 La structure de données de la couche AAL5.

■ La qualité de service dans l'ATM

Conçu à l'origine pour traiter des flux de tout type, ATM différencie les traitements d'extrémité (AAL) et offre à chacun une connexion en adéquation avec ses besoins : c'est la notion de **classe de service**. La classe de service requise est invoquée à la connexion (figure 23.23).

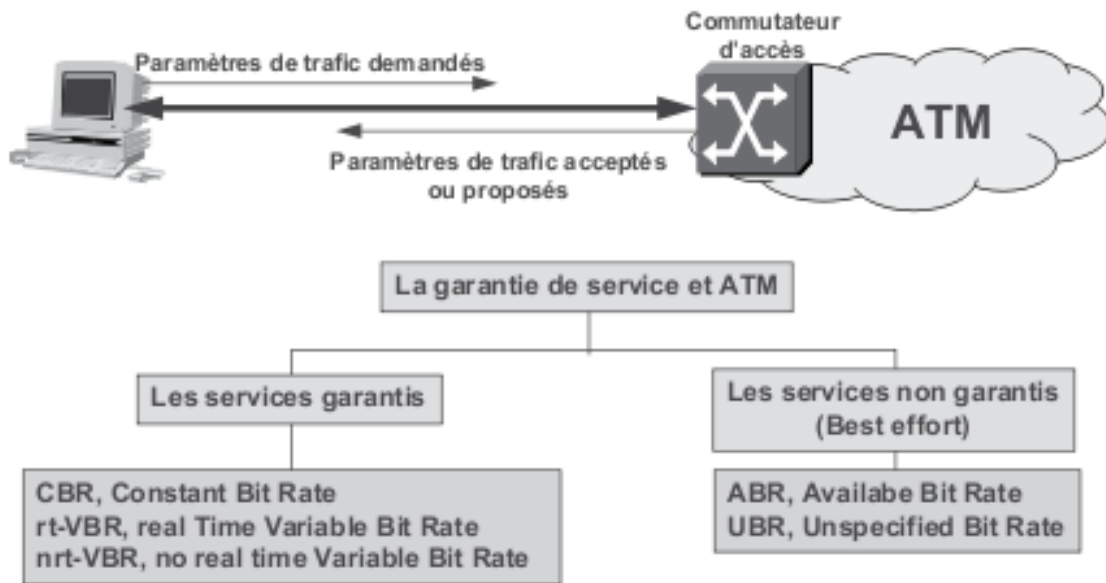


Figure 23.23 Le contrat de service et les classes de service.

La classe de service **CBR** (*Constant Bit Rate* ou **DBR**, *Deterministe Bit Rate*) définit un raccordement à débit constant qui correspond à une émulation de circuit. Elle est destinée aux applications de type voix ou vidéo non compressées.

La classe **VBR** (*Variable Bit Rate* ou **SBR**, *Statistical Bit Rate*) s'applique aux trafics sporadiques. La classe VBR-rt correspond aux applications de type voix ou vidéo compressées. La classe VBR-nrt (*VBR no real time*) est généralement requise pour les applications de type transactionnel.

La classe de service **ABR** (*Available Bit Rate*), dans celle-ci, les applications utilisent le débit disponible sur le réseau (entre les deux bornes prédéfinies). La classe ABR est adaptée à l'interconnexion de réseaux locaux.

De même, une classe de service de type datagramme ou *best effort* a été définie : l'**UBR** (*Unspeciefied Bit Rate*). La classe UBR convient aux applications de type messagerie et sauvegarde à distance (*Remote backup*).

24 MPLS, Multiprotocol Label Switching

Les applications informatiques s'appuient essentiellement sur le protocole TCP/IP. Les données sont généralement transportées par encapsulation dans un protocole en mode connecté (X.25, Frame Relay ou ATM). Indépendamment des problèmes de mise en relation des espaces d'adressage et de la difficulté du respect de la qualité de service de bout en bout, la gestion (établissement, maintien et rupture) des circuits virtuels conduit à des solutions plus ou moins complexes. **MPLS** (*MultiProtocol Label Switching*) pallie cette rupture de technologie en offrant un service en mode connecté transparent aux applications IP. En fait, MPLS migre un réseau IP routé en un réseau IP commuté, il associe la souplesse du routage de niveau 3 à l'efficacité de l'acheminement de niveau 2. C'est un protocole de conciliation généralement présenté comme un protocole de niveau 2-5 (figure 24.1).

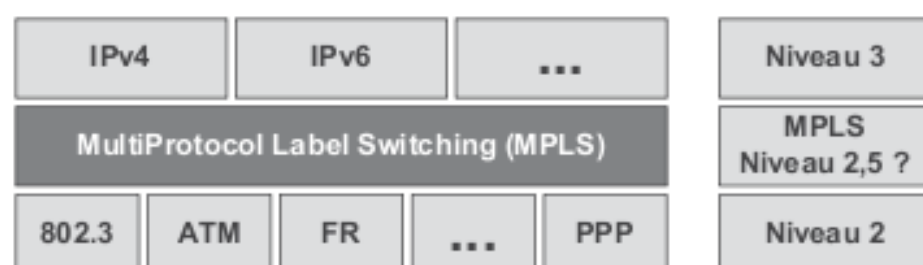


Figure 24.1 Le modèle de référence et l'architecture MPLS.

24.1 Principe

MPLS substitue un protocole de routage IP (RIPv2, OSPF, BGP) au protocole de signalisation propre au réseau de transport. À chaque point de sortie identifié par le protocole de routage, MPLS associe un label (référence de commutation). Chaque datagramme IP se voit, en entrée du réseau MPLS et en fonction de l'adresse destination IP (fonction de routage), attribuer un

label. Par la suite dans le réseau, le datagramme sera acheminé en fonction de ce label (commutation). Un réseau MPLS associe deux fonctions : routage en périphérie et commutation en cœur de réseau (figure 24.2).

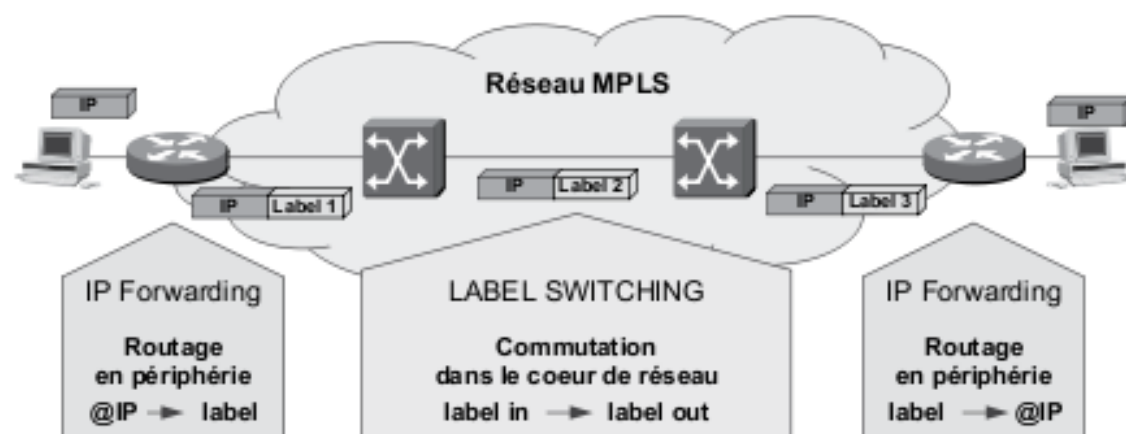


Figure 24.2 Les fonctions de base du réseau MPLS.

24.2 Le réseau MPLS

24.2.1 Le commutateur MPLS

Un réseau MPLS est constitué d'un ensemble de commutateurs MPLS. Un commutateur MPLS ou *Label Switching Router (LSR)* correspond à l'association d'un routeur et d'un commutateur. La figure 24.3 présente la structure d'un tel commutateur. En principe, un commutateur Frame Relay ou ATM peut évoluer en commutateur MPLS par une simple mise à jour du logiciel.

Le plan de routage ou plan de contrôle met en œuvre un protocole de routage IP, un label est associé à chaque adresse IP pouvant être jointe. Un protocole spécifique distribue les labels aux LSR voisins (**LDP**, *Label Distribution Protocol*). Le datagramme IP labellisé en entrée du réseau (paquet MPLS) est ensuite commuté de manière similaire à la commutation traditionnelle (*Port In, Label In* → *Port out, Label Out*). Les labels sont affectés en fonction du nœud de sortie et non de l'adresse IP destination. Tous les paquets à destination d'un même nœud de sortie reçoivent le même label et subissent donc le même traitement dans le réseau. Ce principe, appelé agrégation de routes à la périphérie (figure 24.4), allège les tables d'acheminement et participe à l'amélioration des performances.

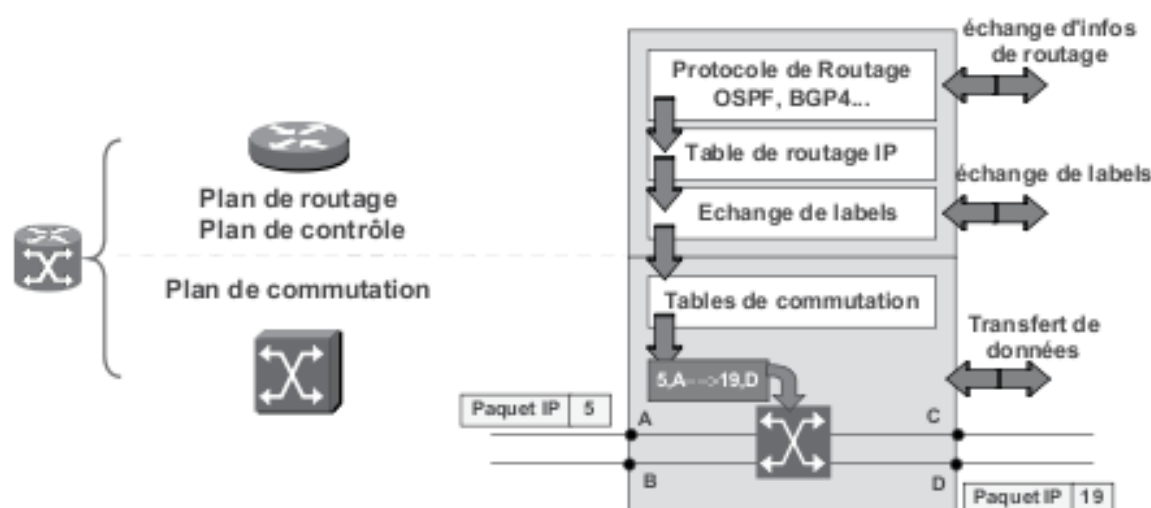


Figure 24.3 La structure d'un Label Switching Router (LSR).

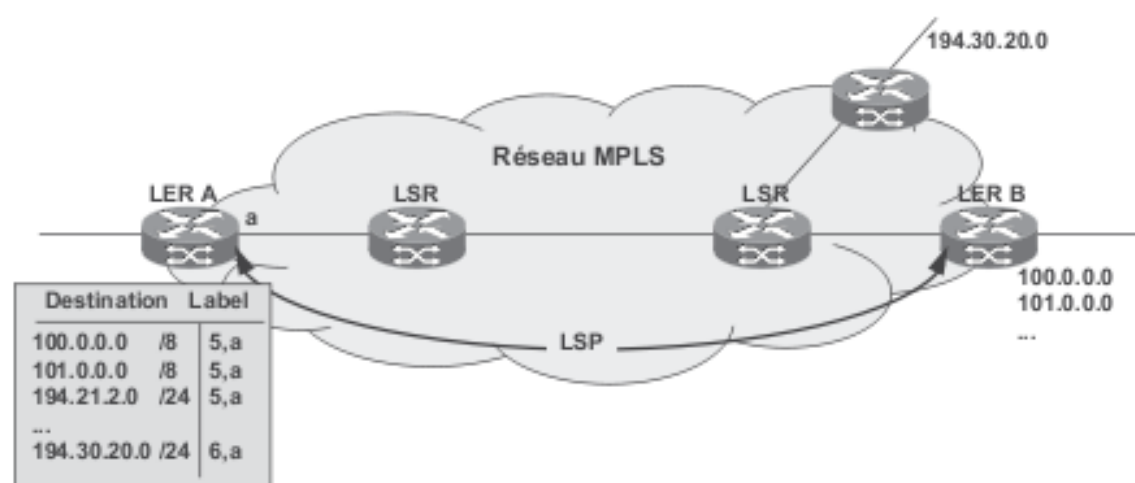


Figure 24.4 Principe de l'agrégation de routes en périphérie.

24.2.2 Principe de l'acheminement dans un réseau MPLS

Dans le réseau de la figure 24.5, toutes les adresses réseaux connues du LER B reçoivent en entrée le même label. L'ensemble des datagrammes qui reçoit un même label forme une **FEC** (*Forwarding Equivalence Class*). Le lien virtuel défini pour une classe d'équivalence est appelé **LSP** (*Label Switched Path*). La figure 8.28 illustre le traitement d'un datagramme IP entre son entrée dans le réseau MPLS et sa sortie.

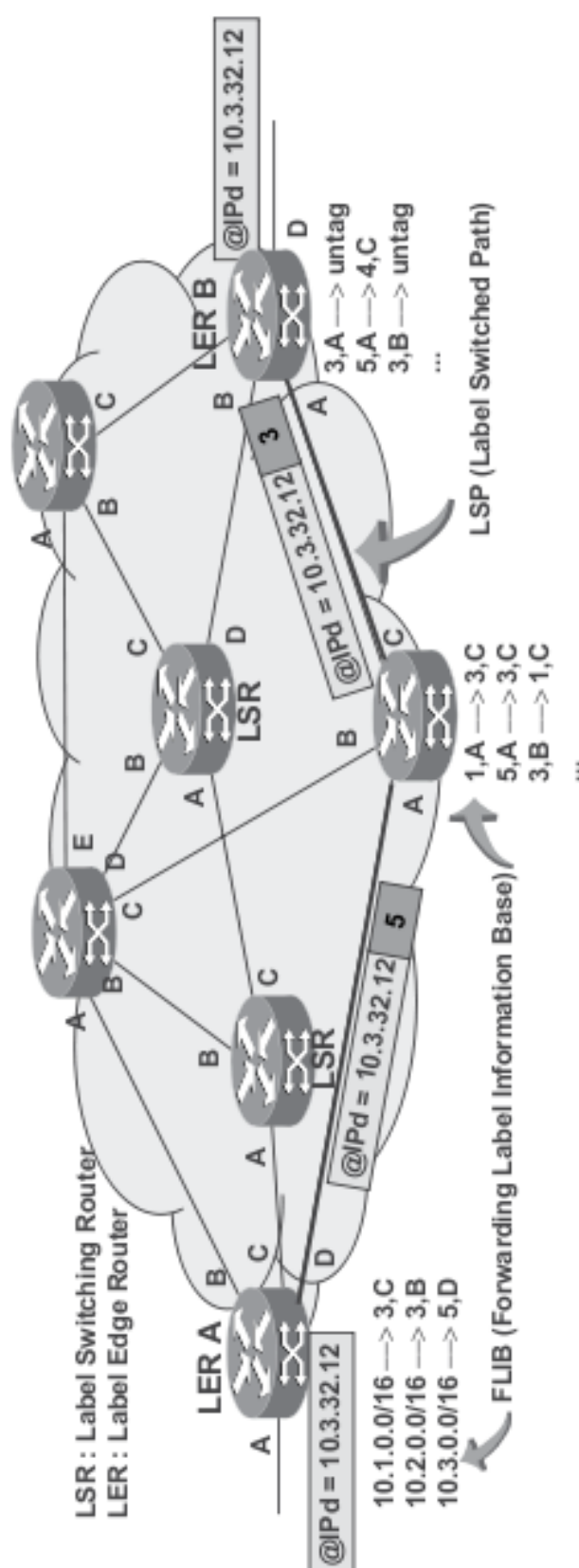


Figure 24.5 L'acheminement d'un datagramme IP dans un réseau MPLS.

Le LER A (*MPLS ingress node*), recevant le datagramme IP consulte sa table de labels (**FLIB**, *Forwarding Label Information Base*). Pour l'adresse IP 10.3.32.12, le label affecté est 5 et le port de sortie sur le LER A est D. Le datagramme est labelisé 5 (*Push tag*). Le LER suivant reçoit sur son port A, un paquet MPLS labelisé 5, il consulte sa table de labels (FLIB), dans laquelle il est indiqué qu'il doit changer le label en 3 et commuter le paquet MPLS en C. Enfin, le LER B de sortie (*MPLS egress node*) examine sa FLIB et constate qu'il doit supprimer le label (*Untag*), il consulte alors sa table de routage IP et achemine le datagramme IP sur le port de sortie D.

Dans un réseau MPLS, un même paquet MPLS peut recevoir plusieurs labels (*Push tag*). L'empilement de labels permet de définir une agrégation de routes en interne dans le réseau et des réseaux virtuels (**VPN**, *Virtual Private Network*). L'opération *Pop tag* permet de supprimer le label de haut de pile, alors que *Untag* supprime le dernier label.

24.2.3 Les mécanismes particuliers

■ L'encapsulation ou l'en-tête MPLS

Hors le champ label sur 20 bits, l'en-tête MPLS ou *shim header* (figure 24.6) permet de gérer différents niveaux de priorité (champ **Exp** pour *Experimental*). Le bit **S** (*Stack*) indique la fin de pile ($S = 1$, dernier label). Le champ **TTL** (*Time To Live*) permet de rendre transparent, à IP, la traversée d'un réseau MPLS. En entrée du réseau, le LER *ingress* recopie la valeur du TTL du datagramme IP dans l'en-tête MPLS, le réseau MPLS décrémente le TTL de la même manière que l'aurait fait un réseau IP. En sortie, le LER *egress* recopie cette valeur dans le datagramme de sortie.

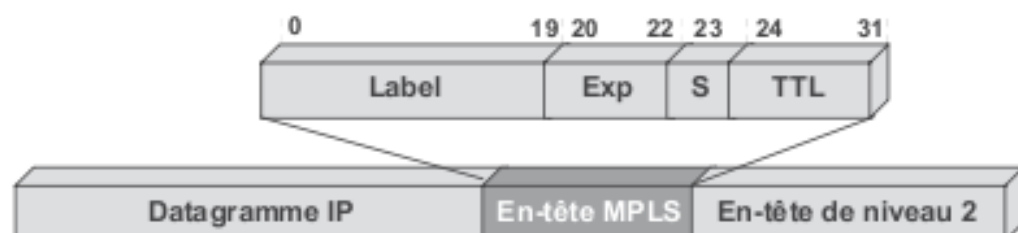


Figure 24.6 Le format de l'en-tête de compensation MPLS.

■ MPLS et les protocoles de niveau 2

La principale difficulté de l'encapsulation de paquets MPLS dans les protocoles de niveau liaison est en relation directe avec la MTU de ces protocoles. La figure 24.7 illustre ces encapsulations. Par exemple, dans le cas du protocole PPP (*Point to Point Protocol*), l'insertion de labels accroît la taille de la trame. Les différents LSR doivent avoir la capacité de fragmenter la trame PPP.

24.3 Les VPN MPLS

24.3.1 Définitions

Un réseau privé est dit virtuel (**VPN**, *Virtual Private Network*) lorsque, sur une infrastructure partagée (réseau public ou privé), on développe des mécanismes tels que la communication ne soit possible qu'entre clients d'un même VPN. Les mécanismes utilisés peuvent être un simple identifiant de VPN (MPLS) ou une technique de chiffrement des communications. Seuls peuvent établir une communication, les clients utilisant le même identifiant ou possédant la même clé de chiffrement/déchiffrement.

24.3.2 Principe général des VPN MPLS

MPLS utilise le modèle *peer* dit aussi *homologue*. Pour assurer la gestion des VPN qui leur sont rattachés, les *Provider Edge Router* entretiennent une table d'acheminement spécifique à chaque interface (**VRF**, *VPN Routing and Forwarding*). La VRF contient la désignation du PE, auquel sont rattachés le client VPN destinataire, la classe d'équivalence associée (**FEC**, *Forwarding Equivalence Class*), considérée par cette table comme le prochain saut (*Next hop*) et le label identifiant le VPN d'appartenance. La VRF contient ainsi tous les éléments d'acheminement, y compris une copie dédiée de la LFIB ; de ce fait, une VRF correspond à un routeur dédié au VPN (**routeur virtuel**). L'acheminement en interne dans le réseau est réalisé par consultation de la **LFIB** (*Label Forwarding Information Base*). La figure 24.8 illustre l'affectation et la commutation de labels d'un VPN MPLS.

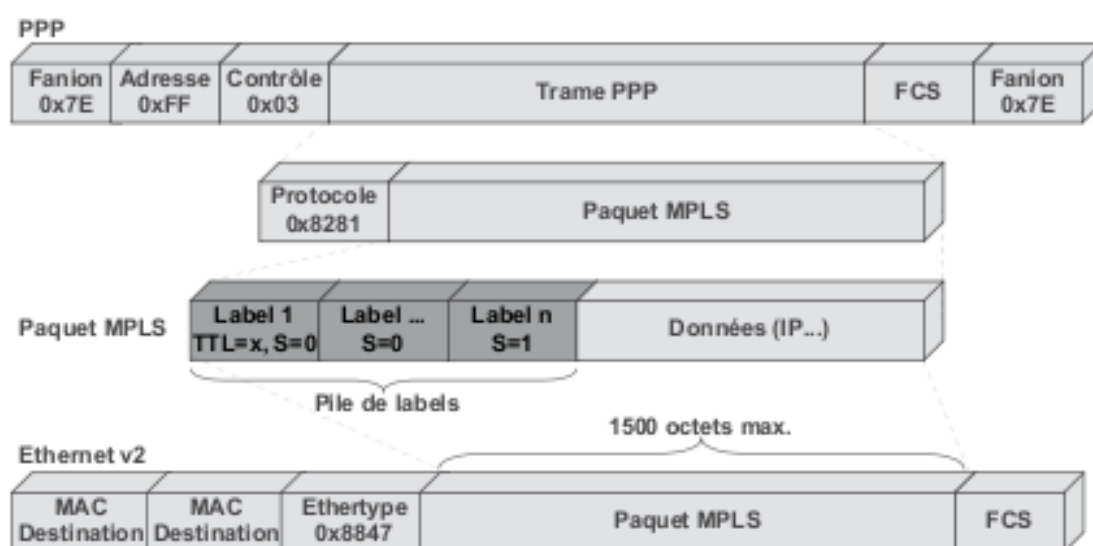


Figure 24.7 MPLS et les protocoles de niveau 2.

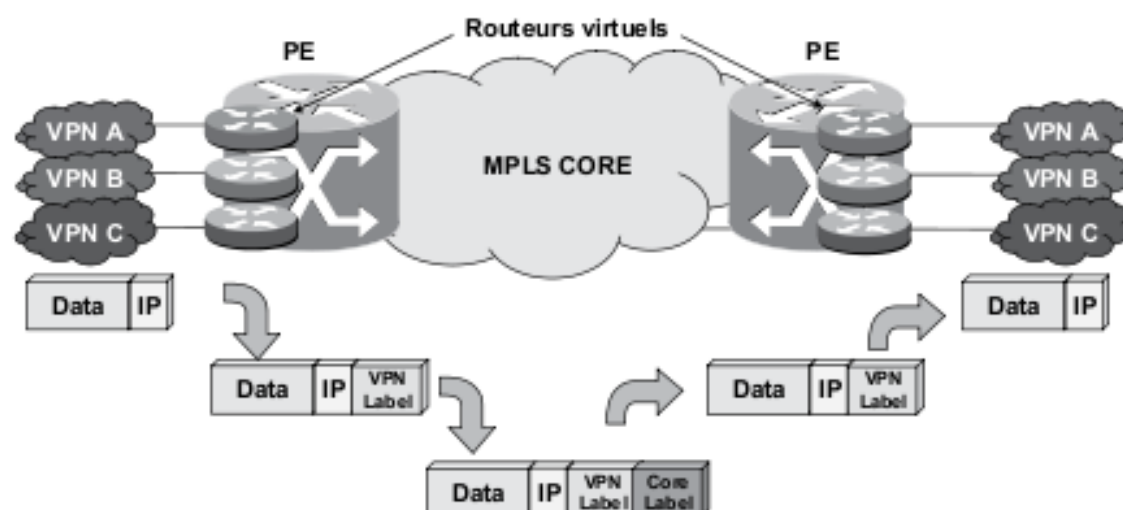


Figure 24.8 L'affectation des labels dans un VPN MPLS.

Dans la figure 24.9, l'entreprise « ENT » dispose d'un VPN entre ses établissements. La VRF du PE indique que tout ce qui provient de l'interface « d » (int d) appartient au VPN de l'entreprise « ENT », que ce VPN est identifié par le label de VPN « ENT » et que dans le VPN « ENT » pour joindre, par exemple, la destination 10.3x.Y.0/24, il faut emprunter la FEC A. Enfin, la FLIB indique que les données destinées à la FEC A sont à remettre à l'interface « a » et que celle-ci est identifiée dans le réseau par le label d'acheminement « La ».

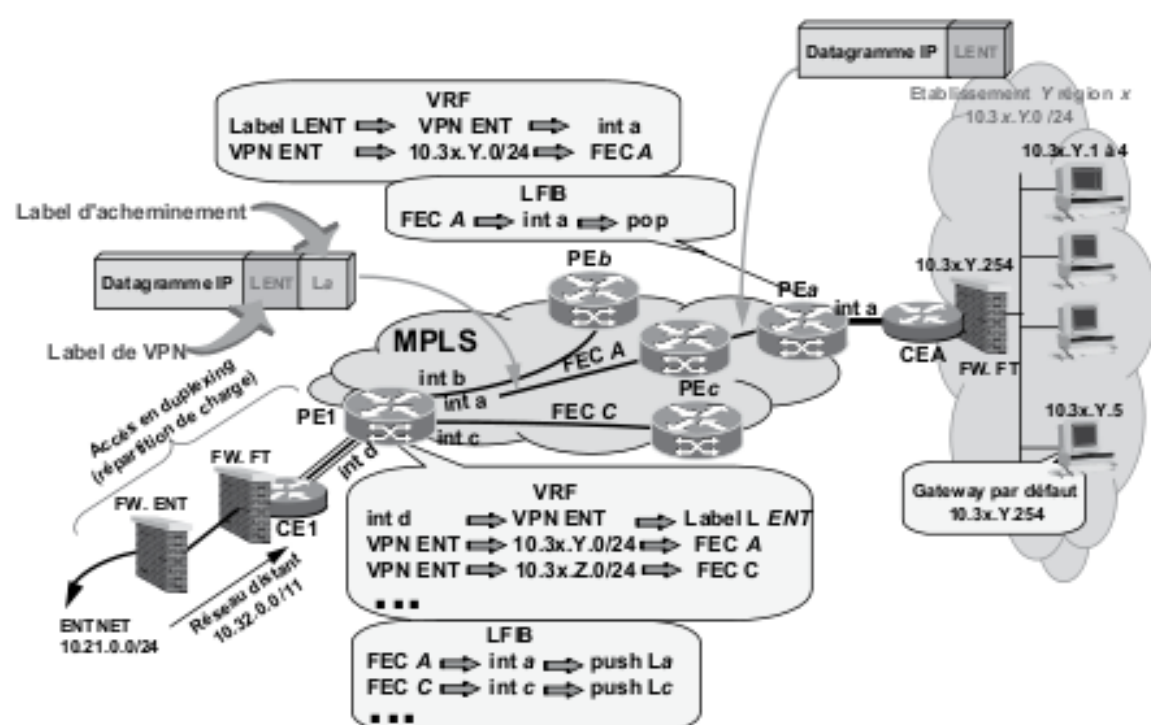


Figure 24.9 L'acheminement dans le réseau MPLS.

25

L'accès aux réseaux, la boucle locale

25.1 Définition

La boucle locale correspond à l'ensemble des moyens mis en œuvre par un opérateur pour collecter le trafic des utilisateurs. Une définition plus restrictive limite l'utilisation du terme boucle locale au seul câble de raccordement usager/réseau.

25.2 Organisation de la distribution des accès

Les moyens d'accès se répartissent en deux catégories, les accès aux réseaux d'opérateur (opérateurs de boucle locale) et les moyens fournis à l'utilisateur pour raccorder ses propres sites informatiques et réaliser ainsi un réseau privé (opérateur de liaisons louées).

La réalisation d'un réseau de distribution (collecte) nécessite des investissements importants. Dans la plupart des pays, ces réseaux ont été financés par des ressources publiques. La mise en concurrence des télécommunications a donc posé le problème du partage de cette ressource. C'est sous le terme de dégroupage de la boucle locale que l'ARCEP (Autorité de régulation des communications électroniques et des postes) a organisé ce partage en instituant la colocalisation des équipements actifs (figure 25.1).

La notion de dégroupage se décline selon 2 modes accès à la boucle locale. Le dégroupage total donne, à un opérateur autre que France Télécom, l'accès à toute la bande de fréquences de la paire de cuivre. L'ensemble

des services, téléphonie et accès Internet est alors complètement géré par l'opérateur alternatif. Dans le dégroupage partiel, l'opérateur alternatif n'a accès qu'aux fréquences hautes de la paire de cuivre, les services téléphoniques (fréquences basses) restent gérés par France Télécom.

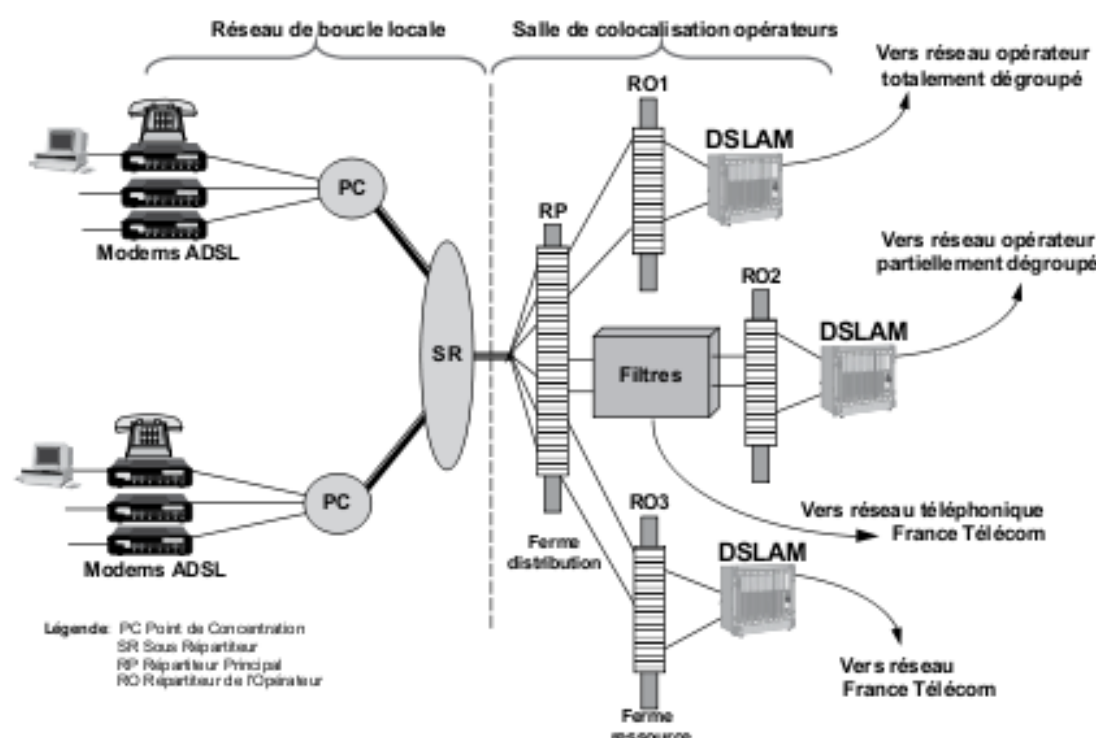


Figure 25.1 Le dégroupage de la boucle locale et l'accès ADSL.

25.3 Les accès haut débit

Dans les grandes métropoles, la diffusion de programmes vidéo a largement fait appel à la fibre optique qui constitue le support idéal pour les transmissions large bande (**PON**, *Passive Optical Network*). Amener la fibre optique chez l'abonné (**FTTH**, *Fiber To The Home*) est une solution idéale mais coûteuse (figure 25.2). Aussi, la distribution chez l'abonné final a été généralement réalisée par un câble cuivre à partir d'un point de distribution commun situé dans l'immeuble (**FTTB**, *Fiber To The Basement*) ou au plus près d'un groupe d'habitations (**FTTC**, *Fiber To The Curve*). Cependant, le développement d'Internet et le

besoin en bande passante qu'il a révélé conduisent aujourd'hui les opérateurs français à rattraper leur retard et à proposer des accès de type FTTH à 100 Mbit/s.

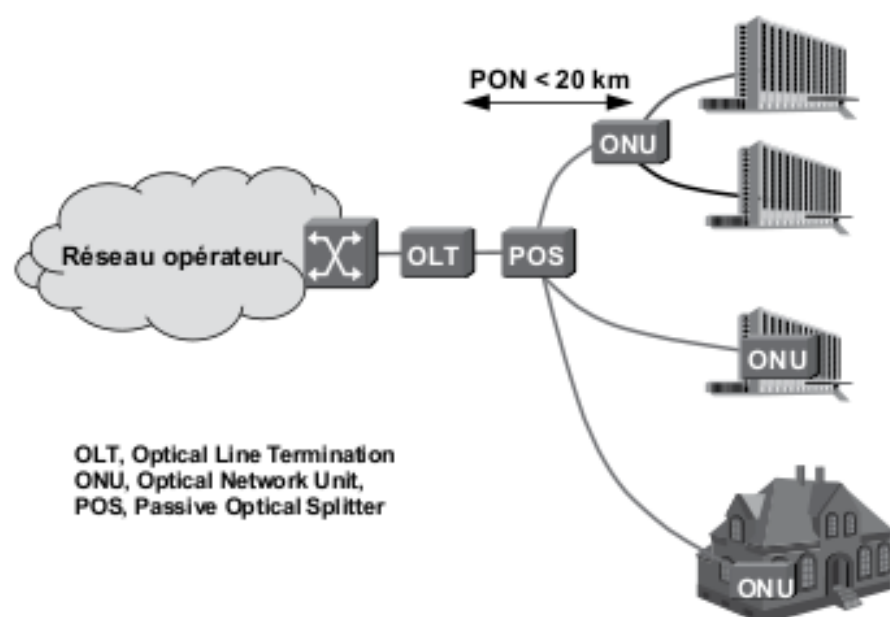


Figure 25.2 Principe de la boucle locale optique.

25.3.1 Les techniques DSL

La bande passante du service voix est limitée à 4 kHz, cependant la bande passante réelle de la paire torsadée est supérieure au MHz. La technologie **ADSL** (*Asymmetric data rate Digital Subscriber Line*) partage la bande disponible entre le service voix analogique (0 à 4 kHz) et deux canaux de données simplex, l'un dit montant (*Up*) offre une bande passante de 32 à 640 kbit/s, le second (*Down*) un débit de 1,5 à 8,2 Mbit/s.

Les données sont transposées en fréquence selon un codage spécifique dit **DMT**¹ (*Discrete MultiTone*) dans la bande de 25 kHz à 1,1 MHz². Le codage DMT divise chacun des spectres haut débit en sous-canaux (por-

1 La modulation DMT est aussi appelée OFDM (*Orthogonal Frequency Division Multiplexing*).

2 En utilisant une bande de fréquences étendue jusqu'à 2,2 MHz et 512 porteuses, l'ADSL2+ autorise un débit théorique pouvant atteindre 25 Mbit/s et un débit montant de 1 Mbit/s.

teuses ou tonalités) espacés de 4,312 kHz. Chaque sous-canal, modulé en phase et en amplitude (MAQ) codant, en principe, 8 bits dans un temps d'horloge, constitue un symbole DMT. En adaptant le nombre de bits par symbole et le nombre de sous-canaux utilisés en fonction de l'état de la ligne, la technique DMT optimise en permanence le débit en fonction de la qualité du canal de transmission par pas de 32 kbit/s (figure 25.3).

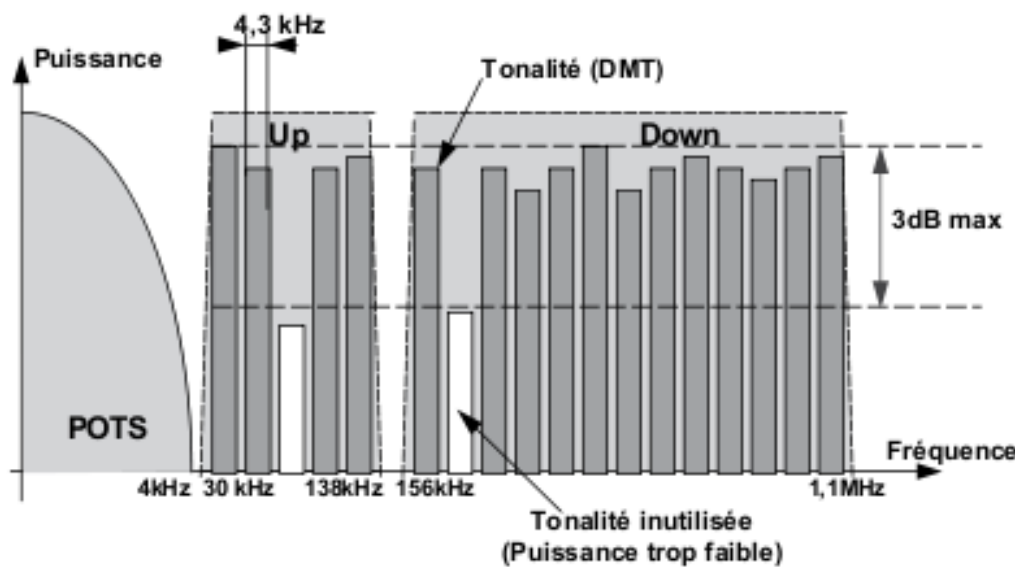


Figure 25.3 Le spectre DMT utilisé dans l'ADSL.

L'accès au réseau haut débit de l'opérateur *via* la ligne téléphonique nécessite l'installation d'un équipement spécifique chez l'utilisateur final qui assure la séparation des canaux : le *splitter* (séparateur vocal), ou coupleur **POTS** (*Plain Old Telephon Service*, service téléphonique traditionnel) et le modem ADSL (x_Box). Le *splitter* est généralement intégré au modem. Le modem ADSL offre un accès de type Ethernet, USB ou Wi-Fi. Du côté opérateur, le **DSLAM** (*Digital Subscriber Line Access Multiplex*) est un multiplexeur fréquentiel assurant la séparation des bandes de fréquences téléphoniques et de données, c'est aussi un modem ADSL qui dans cette fonction assure l'interface entre les connexions utilisateurs et le réseau haut débit de l'opérateur (figure 25.4). Cependant, la connexion IP doit être prolongée jusqu'à chez le fournisseur de service Internet (**FAI**), ce rôle est dévolu au **BAS** (*Broadband Access Server*).

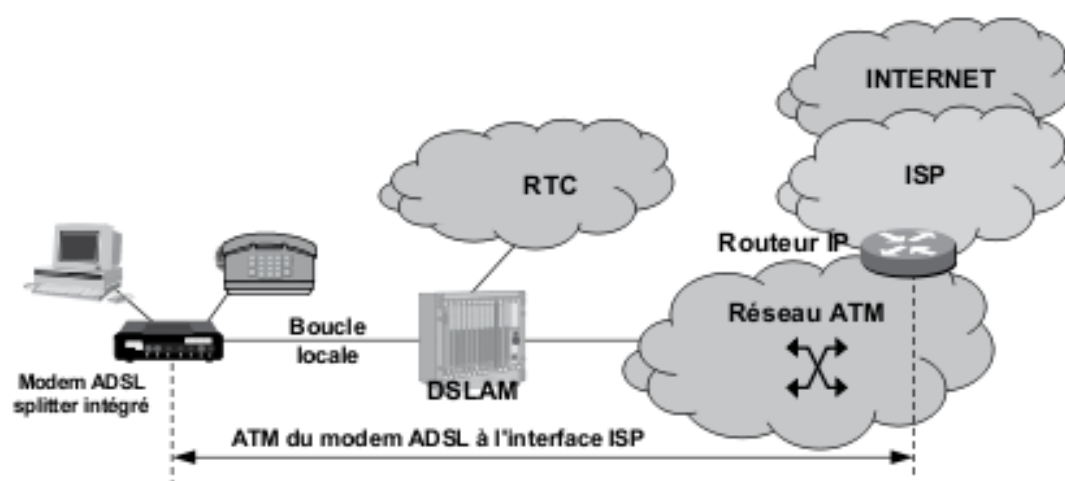


Figure 25.4 L'architecture d'un réseau ADSL.

C'est le BAS qui réalise les fonctions d'authentification du client et l'acheminement des données vers les différents fournisseurs d'accès. Cette authentification est réalisée par le protocole PPP adapté à une interface Ethernet (**PPPoE**, RFC 2516 *PPP over Ethernet*) ou à une interface USB (**PPPoA**, RFC 2684 *PPP over ATM*). La figure 8.30 (chapitre 8) illustre l'architecture protocolaire d'une liaison ADSL.

Compte tenu de l'intérêt économique des techniques DSL, d'autres solutions ont été développées pour raccorder à moindre frais les usagers aux réseaux des opérateurs. Le tableau 25.1 présente succinctement ces différentes versions.

Tableau 25.1 Les différentes technologies xDSL.

Appellation	Débit descendant	Débit montant	Distance	Utilisation
ADSL	32 kbit/s à 8 Mbit/s	32 kbit/s à 1,1 Mbit/s	5,5 km	Accès professionnel à Internet Interconnexion de LAN Vidéo à la demande (VoD)
UADSL G.Lite	64 kbit/s à 1,5 Mbit/s	32 kbit/s à 512 kbit/s	5,5 km	Accès résidentiel à Internet (obsolète)
SDSL	1,54 Mbit/s (T1) 2,048 Mbit/s (E1)	1,54 Mbit/s (T1) 2,048 Mbit/s (E1)	6,5 km	Interconnexion de LAN Serveur Internet Vidéoconférence

Appellation	Débit descendant	Débit montant	Distance	Utilisation
IDSL	144 kbit/s à 1,5 Mbit/s	32 kbit/s à 512 kbit/s	11 km	Accès RNIS
VDSL	13 à 52 Mbit/s	1,5 à 2,3 Mbit/s	1,2 km	Accès Internet, VoD TV haute définition
HDSL	1,5 Mbit/s (T1) 2,048 Mbit/s (E1)	1,5 Mbit/s (T1) 2,048 Mbit/s (E1)	4,5 km	Accès professionnel E1 Raccordement PABX Interconnexion de LAN

26 Ethernet dans les MAN et WAN

26.1 Les réseaux sans coupure

Si MPLS apporte à l'interconnexion des réseaux un traitement IP de bout en bout, il n'en demeure pas moins qu'au niveau transport nous avons encore des ruptures de technologie entre l'accès local en technique xDSL, le transport dans le WAN en Frame Relais ou ATM sur SDH. Aussi, le *Metropolitan Ethernet Forum* (MEF) propose de substituer aux réseaux traditionnels un transport Ethernet de bout en bout (MEN, *Metropolitan Ethernet Network*) dont le principe est illustré par la figure 26.1.

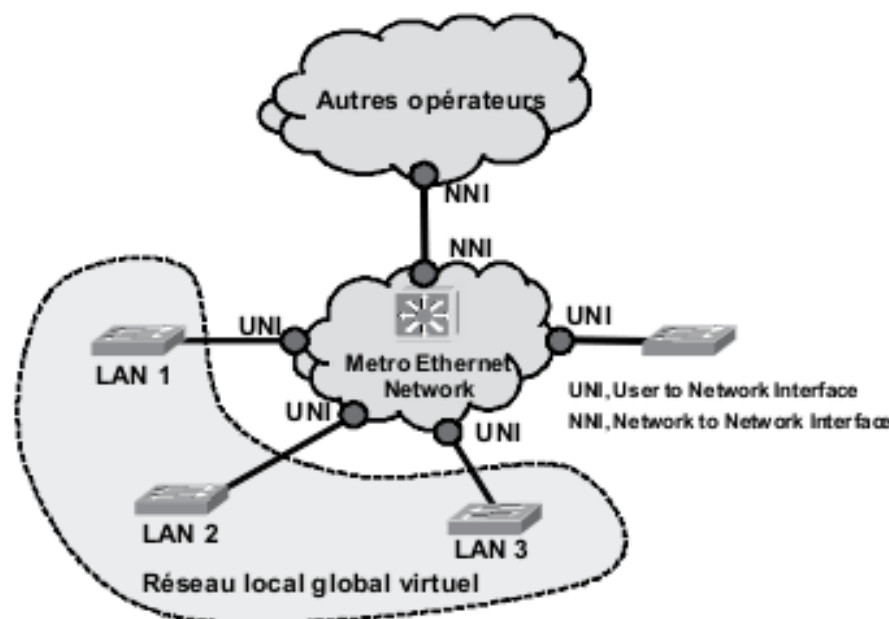


Figure 26.1 Principe d'un réseau sans couture.

Dans la figure 26.1, le MEN est vu par les différents réseaux locaux de l'entreprise comme un commutateur *backbone*, ainsi, les différents LAN ne constituent qu'un seul LAN virtuel.

26.2 Ethernet à grande distance (CGE, Carrier Grade Ethernet)

Dans les LAN, Ethernet commuté achemine les trames par consultation d'une table construite par apprentissage, toute trame dont l'adresse n'est pas connue est acheminée par diffusion (fonctionnement HUB). Ce mode de fonctionnement est inconcevable dans une infrastructure publique partagée ; aussi, Ethernet grande distance doit-il migrer un service sans connexion avec un apprentissage distribué des adresses vers un service de multiples tunnels orientés connexion (EVC, *Ethernet Virtual Connection*) mis en œuvre par un système centralisé d'acheminement. De même, MEN (*Metro Ethernet Network*) étant une infrastructure publique, doit admettre divers modes d'accès (figure 26.2).

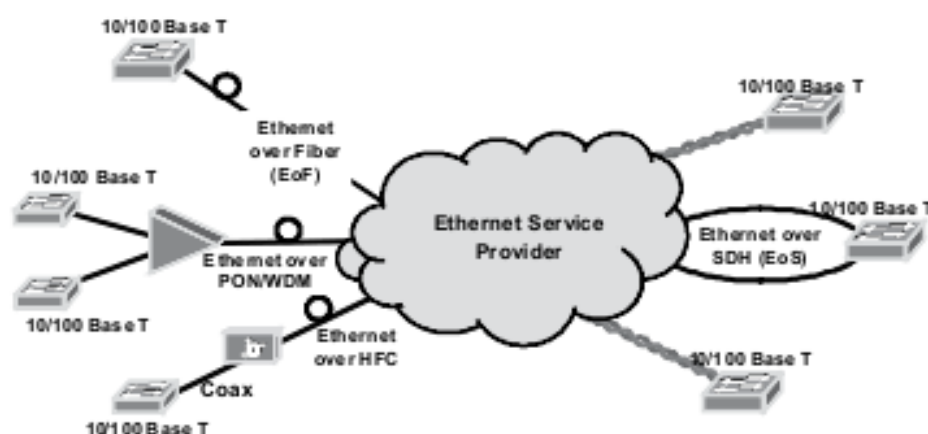


Figure 26.2 Ethernet Carrier Grade, les modes d'accès.

Migrer d'un environnement local à un environnement opérateur impose de nombreuses contraintes, notamment de définir des équipements dits de « classe opérateur » en termes de fiabilité et de service (Résilience, QoS, redondance, supervision).

La résilience consiste non seulement à garantir la fiabilité des équipements¹, mais surtout à détecter les pannes et à restaurer un service acceptable dans un temps bref (50 ms pour SDH).

1 Un équipement de classe opérateur doit avoir un taux de disponibilité important, en téléphonie ce taux répond à la règle des cinq 9 (99,999%), soit une indisponibilité inférieure à 5 minutes/an.

La Qualité de service de la recommandation IEEE 802.1p ne permet que 8 niveaux de priorité (8 bits) alors que l'IETF en définit 14 dans DiffServ. Le tableau 26.1 indique la relation proposée par la MEF entre les services DiffServ et IEEE 802.1p.

Tableau 26.1 La QoS dans Ethernet Carrier Grade.

Valeur du champ CoS	Correspondance DiffServ
2	DiffServ Best Effort (BE 00)
3	DiffServ Assured Forwarding, probabilité de perte importante (AF 33)
4	DiffServ Assured Forwarding (AF 32)
5	DiffServ Assured Forwarding, probabilité de perte faible (AF 31)
6	DiffServ Expedited Forwarding (EF 46)

Le Metro Ethernet Forum (MEF) a défini trois types de services (figure 26.3) pour les réseaux CGE (MEF 6-1, *Metro Ethernet Services Definitions Phase 2*) :

- ▶ les services *E-Ligne* réalisent une connexion sécurisée entre deux sites (mode point-à-point) ;
- ▶ les services *E-LAN* qui assurent une interconnexion totale de plusieurs sites ;
- ▶ les services *E-Tree* qui prennent en charge les flux à destination multiple comme la diffusion d'émission de télévision sur IP (IPTV).

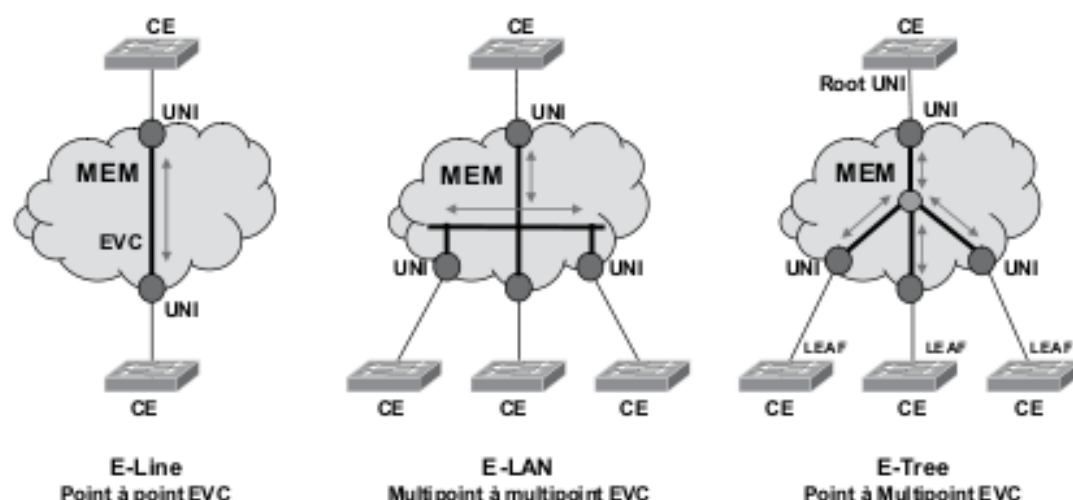


Figure 26.3 Les services du MEF.

Dans les LAN, la détection de collision réalise un contrôle flux natif. En mode *full-duplex* la détection de collision étant invalidée, l'IEEE a introduit un mécanisme de type ON-OFF par l'émission de trames « Pause » (IEEE 802.3x). Le transport de la ToIP sur les MAN Ethernet interdit toute notion de contrôle de flux, aussi afin de garantir une certaine qualité de service, une connexion ne peut être acceptée que si le réseau est apte à la satisfaire. À cette fin, le MEF (MEF 23, *Class of Service Implementation Agreement Part 1*) a introduit un mécanisme de SLA (*Service Level Agreement*) similaire à celui mis en œuvre dans les réseaux Frame Relay¹, toute connexion est définie selon ces paramètres :

- ▶ CIR (*Committed Information Rate*) ;
- ▶ CBS (*Committed Burst Size*) ;
- ▶ PIR (*Peak Information Rate*) ;
- ▶ MBS (*Maximum Burst Size*).

26.3 Modèle architectural

Pour assurer un service Ethernet de bout en bout, le MEF n'a pas redéfini la trame Ethernet, les réseaux MEN utilisent le format traditionnel marqué IEEE 802.1pQ. Le modèle architectural retenu est inspiré du modèle retenu par UIT pour les réseaux RNIS, c'est un modèle en plan représenté figure 26.4.

¹ Voir le paragraphe 2.2.3 Introduction au Frame Relay de ce même chapitre.

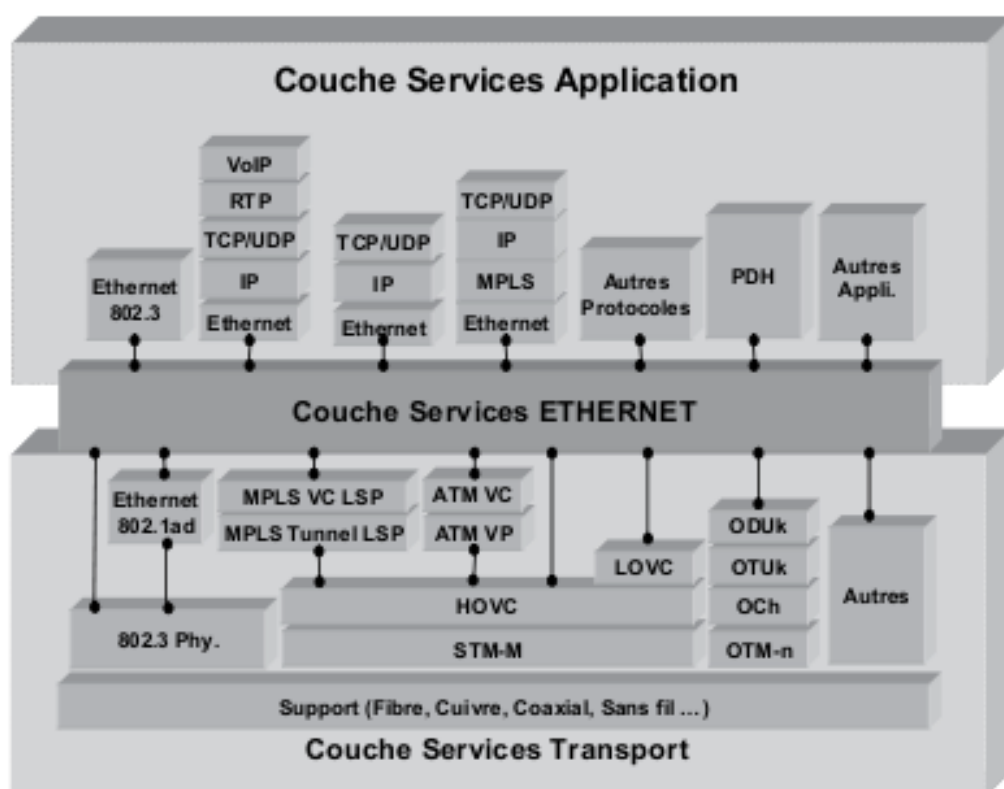


Figure 26.4 Modèle architectural du CGE.

27

Sécurisation des accès

La permanence de l'accès au système d'information est devenue un enjeu stratégique pour les entreprises. Aussi, de nombreuses mesures sont prises pour garantir la sécurisation des liens vers l'extérieur. La figure 27.1 illustre les différents moyens mis en œuvre pour raccorder un réseau d'entreprise au réseau de transport de l'opérateur, en partant du moins sécurisé (lien unique), au plus sécurisé (double adduction sur deux points d'accès différenciés).

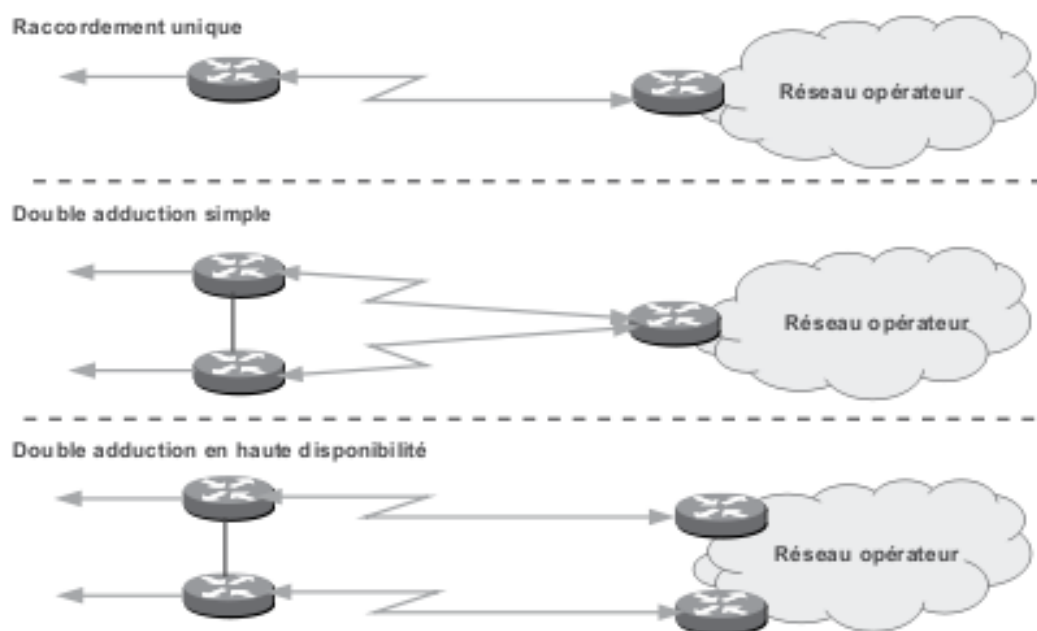


Figure 27.1 Principe d'accès aux réseaux d'opérateur.

Le protocole **VRRP** (*Virtual Router Redundancy Protocol*, RFC 2338) assure la permanence du service d'accès. En cas de rupture d'un lien, le système bascule automatiquement sur le lien de secours. VRRP permet à un ensemble de routeurs situé dans un même domaine de diffusion

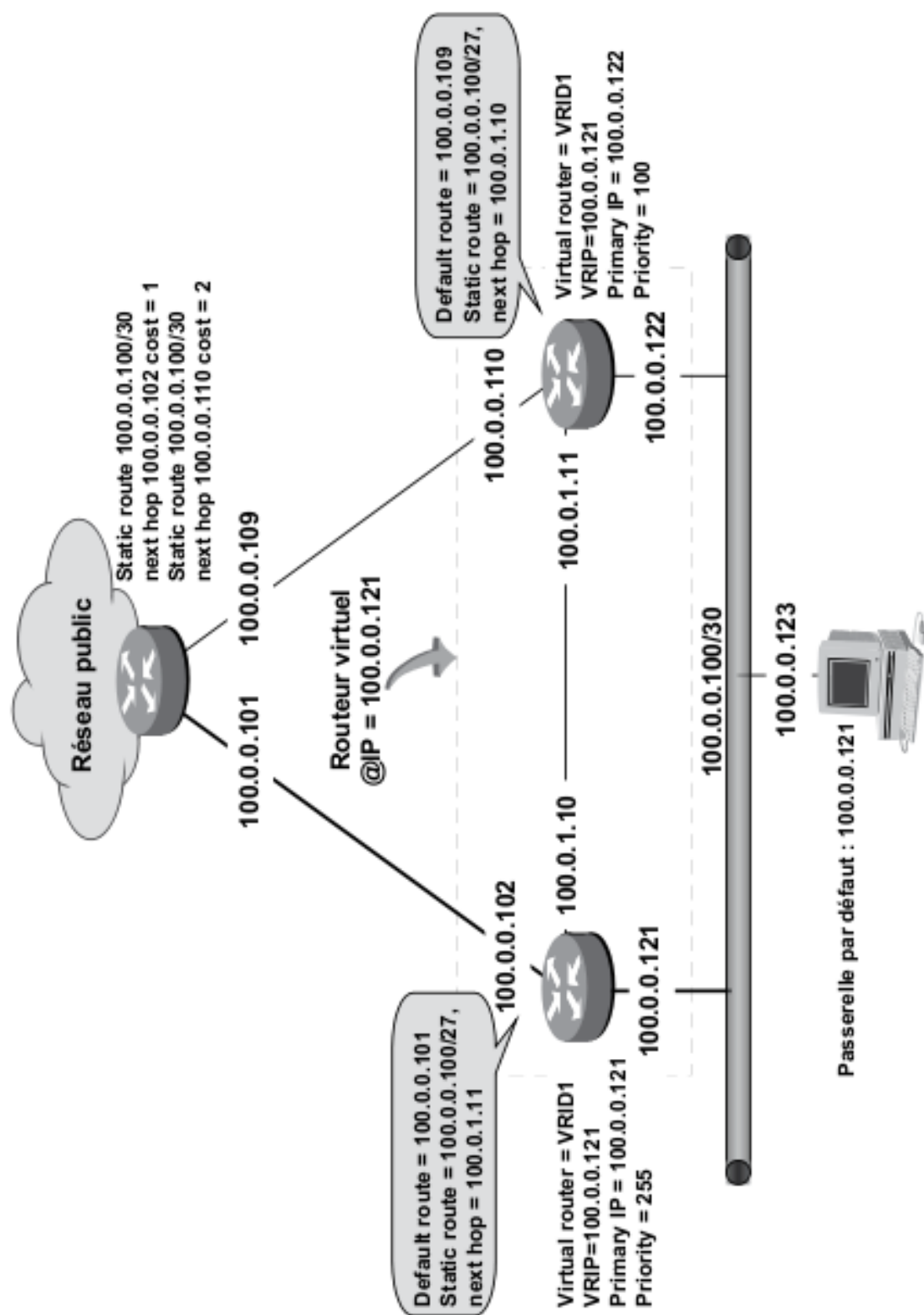


Figure 27.2 Principe du protocole VRRP.

d'être vu par une seule et même adresse IP (IP virtuelle) et une seule et même adresse MAC (MAC virtuelle). Un seul routeur est actif à la fois. La figure 27.2 illustre le principe du protocole VRRP.

Différentes priorités sont affectées aux routeurs du groupe VRRP qui déterminent, en cas de défaillance du routeur principal, l'ordre d'élection du routeur actif. Le routeur actif par défaut se voit attribuer la plus haute priorité, soit 255. En permanence, le système émet des trames de contrôle VRRP (*Hello*) à l'adresse *multicast* 224.0.0.18 et à l'adresse MAC associée 01:00:5E:00:00:xx où xx représente l'identifiant du groupe de routeur (VRID). La défaillance du routeur actif est détectée par l'arrêt de ces diffusions. La figure 27.3 illustre les différents cas de dysfonctionnement.

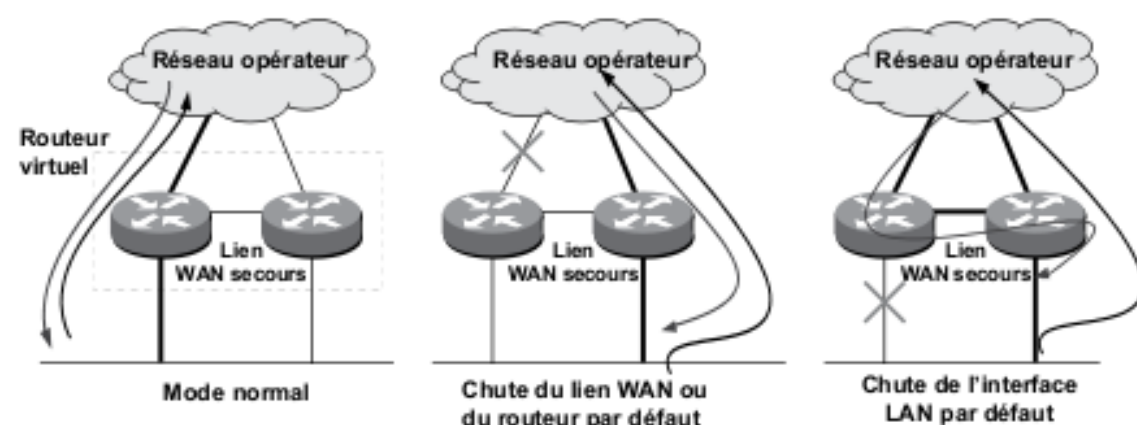


Figure 27.3 Principe du protocole VRRP.

Le protocole VRRP fonctionne en mode Actif/Passif ou Actif/Actif en répartition de charge statique ; des implémentations propriétaires comme HSRP de Cisco (*Hot Standby Routing Protocol*) autorisent un mode Actif/Actif en répartition de charge dynamique par session. Un routeur, dit routeur maître, distribue les flux de manière équitable entre tous les routeurs du domaine HSRP.

27.1 Conclusion

MPLS apporte au protocole IP, non orienté connexion, les avantages du mode connecté tout en conservant la souplesse du mode non connecté. En outre il permet :

- ▶ la diminution du temps de transit dans les réseaux, concurrencé aujourd'hui par les giga- voire téra-routeurs ;
- ▶ l'ingénierie de trafic, qui autorise dans un réseau IP une haute disponibilité et la prévention de la congestion ;
- ▶ la mise en œuvre de la QoS qui a autorisé un service voix de qualité sur un réseau IP ;
- ▶ une gestion souple des VPN.

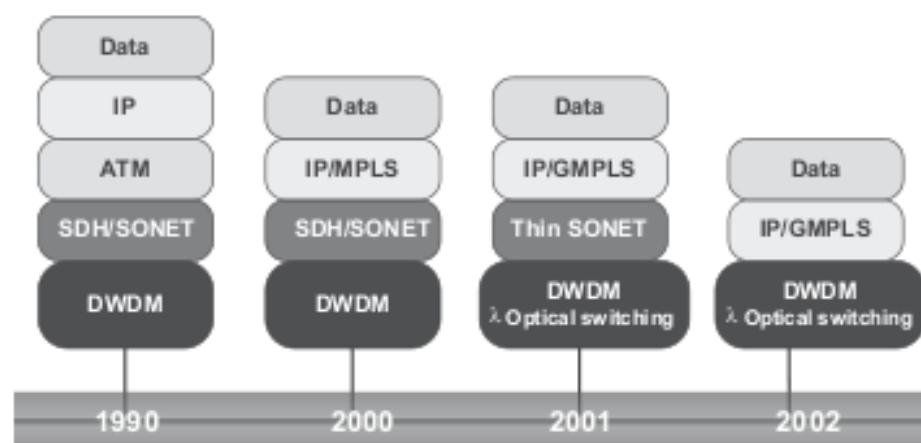


Figure 27.4 Évolution des empilements protocolaires.

Dans l'état actuel des techniques, les réseaux sont constitués d'un empilement de protocoles ayant chacun leur protocole de signalisation et d'administration ainsi que des techniques de configuration spécifiques (figure 27.4). Cependant, si MPLS inséré entre IP et le protocole de transport de niveau 2 optimise les performances du réseau et simplifie la gestion de la bande passante, pour optimiser l'utilisation de celle-ci et faciliter l'administration globale des réseaux, il conviendrait de redéfinir une architecture globale. C'est l'objectif de **G-MPLS** (*Generalized MPLS*) qui étend le domaine MPLS jusqu'au niveau optique en faisant correspondre un label à une longueur d'onde (**MPλS**).

Entre les services de MPLS de niveau 3 et ceux d'*Ethernet Carrier Grade* de niveau 2, il n'y a pas concurrence mais complémentarité. MPLS apporte la souplesse de création de réseaux virtuels et *Ethernet Carrier Grade* l'efficacité d'un acheminement de niveau 2. Ces deux protocoles cohabiteront dans une infrastructure Ethernet/MPLS.

9

Interconnexion des réseaux et la qualité de service

28

L'interconnexion des réseaux

28.1 Définition

Le déploiement des réseaux d'établissement a permis le traitement local des informations. Cependant, pour assurer la cohérence du système d'information de l'entreprise, il est nécessaire d'organiser des échanges d'information entre ses différentes composantes. Tel est l'objet de l'interconnexion des réseaux. Indépendamment de la distance qui les sépare et des protocoles utilisés, l'interconnexion consiste à mettre en relation des machines appartenant à des réseaux physiquement distincts par l'intermédiaire d'éléments d'interconnexion, appelés relais dans la terminologie OSI. Le relais peut n'être qu'un simple élément matériel (pont, routeur...) mais aussi un réseau (figure 28.1).



Figure 28.1 Principe de l'interconnexion des réseaux.

28.2 Les problèmes liés à l'interconnexion

La mise en relation de deux systèmes peut se réaliser très simplement si ces deux systèmes et le relais utilisent les mêmes protocoles, comme, par exemple, l'interconnexion de deux réseaux locaux utilisant TCP/IP *via* un réseau IP comme l'illustre la figure 28.2. Les données issues d'un réseau

sont transportées par un protocole appartenant à la pile TCP/IP : **PPP** (*Point-to-Point Protocol*).

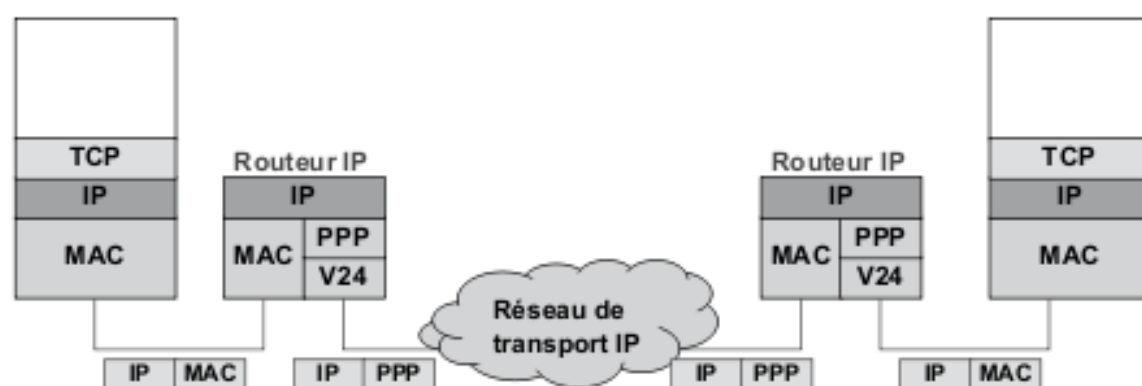


Figure 28.2 L'interconnexion de réseaux homogènes.

Cependant, dans la plupart des cas, le protocole du réseau relais est différent du protocole local, par exemple l'interconnexion de deux réseaux locaux TCP/IP *via* un réseau de transport X.25, Frame Relay (relais de trames) ou ATM. Quand les deux éléments à raccorder mettent en œuvre des technologies différentes, l'hétérogénéité est de bout en bout. Dans ce cas, pour assurer l'interfonctionnement des systèmes, une unité d'interfonctionnement (UIF) réalise les adaptations nécessaires. Trois techniques peuvent alors être utilisées : la conversion de service, la conversion de protocole et l'encapsulation.

28.3 L'encapsulation ou *tunneling*

La conversion de protocole est irréversible. En effet, si on interconnecte deux réseaux IP *via*, par exemple, un réseau X.25, il est aisé, à partir des informations d'en-tête du datagramme IP de confectionner un en-tête X.25, mais il est totalement impossible à partir des données d'en-tête X.25 de reconstituer le datagramme d'origine. Pour disposer des informations nécessaires à la reconstruction du datagramme d'origine, il conviendrait d'insérer un sous-champ entre l'en-tête du nouveau protocole et les données à transporter contenant les informations manquantes. Cette méthode serait lourde, il est préférable de transporter, dans le champ de données d'X.25, le datagramme IP complet (données et en-tête). Ce mécanisme se

nomme **encapsulation de données**. On parle aussi de *tunneling* car on a réalisé un « tunnel » X.25 transportant des données IP. L'encapsulation de données est illustrée par la figure 28.3. Les données sont transférées en mode transparent. Le relais final exécute l'opération inverse, l'unité de données du protocole d'origine est restituée.

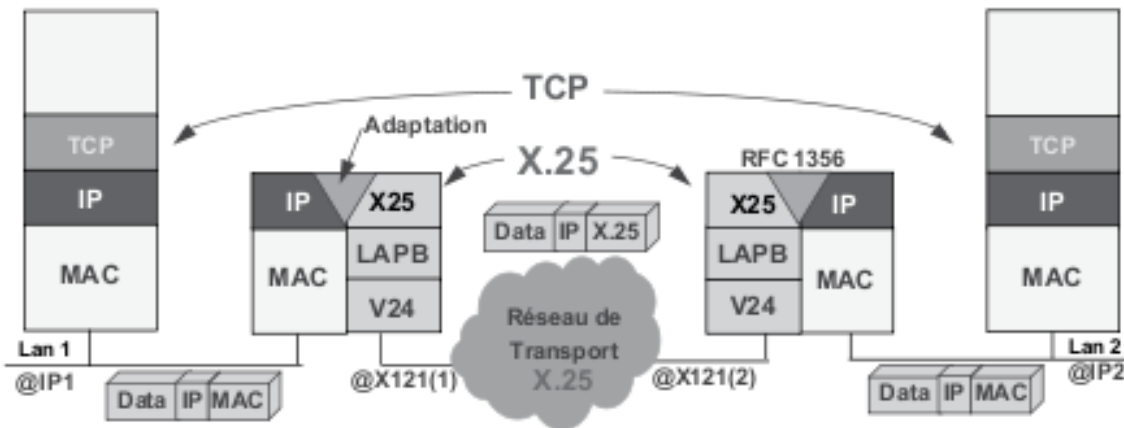


Figure 28.3 L'interconnexion de réseaux hétérogènes.

Les deux réseaux locaux IP LAN 1 et LAN 2 sont reliés *via* un réseau X.25 (RFC 1356). Le protocole inter-réseau utilisé (X.25) est totalement incompatible avec le protocole des réseaux d'extrémité (TCP/IP). Non seulement les adresses sont incompatibles (X.121 pour X.25 et IP pour TCP/IP), mais les réseaux utilisent des techniques différentes (mode datagramme pour TCP/IP et mode orienté connexion pour X.25). La passerelle devra assurer la conversion d'adresses, l'ouverture et la fermeture sur temporisation des circuits virtuels et ceci pour chaque adresse IP distante. Le datagramme IP est transporté de manière transparente par le protocole X.25.

29 Les éléments d'interconnexion (relais)

29.1 Définitions

Selon le niveau du modèle de référence où se réalise l'interconnexion de deux ou plusieurs réseaux, l'ISO distingue quatre types de relais (figure 29.1) :

- ▶ les **répéteurs**, organes d'interconnexion locaux, agissent au niveau 1 du modèle de référence. Un répéteur ne fait que retransmettre d'un côté les bits reçus de l'autre, il agit par diffusion, il peut être utilisé pour prolonger un support, changer de support en encore réaliser l'isolation galvanique entre deux segments d'un réseau ;
- ▶ les **ponts** (*bridges*) sont des éléments d'interconnexion de niveau 2. Ils permettent d'interconnecter deux ou plusieurs réseaux (ponts multi-ports) dont les couches physiques sont semblables ou dissemblables. Les ponts sont transparents aux protocoles de niveau supérieur. Les ponts assurent des fonctions d'adaptation de débit ou de support entre réseaux locaux. Agissant au niveau 2 du modèle de référence, ils ont accès à l'adresse MAC. De ce fait, ils peuvent acheminer des trames en fonction de l'adresse MAC réalisant ainsi un acheminement de niveau 2. Les ponts sont aujourd'hui remplacés par des commutateurs qui ne sont que des ponts améliorés ;
- ▶ les **routeurs** sont des éléments d'interconnexion de niveau 3 qui acheminent (routent) les données vers un destinataire identifié par son adresse de niveau 3 (IP du DoD ou autre). Les routeurs offrent plus de possibilités que les ponts puisqu'ils peuvent mettre en œuvre les mécanismes du niveau 3 (segmentation, réassemblage, contrôle de congestion...) ;

- au-dessus, on parle de **passerelles**. Généralement les passerelles agissent au niveau applicatif, elles mettent en relation des applications de même nature mais d'architecture différente ; c'est, par exemple, le cas d'une interconnexion d'une messagerie **SMTP** (*Simple Mail Transfer Protocol*) du monde TCP/IP avec une messagerie X.400 du monde ISO.

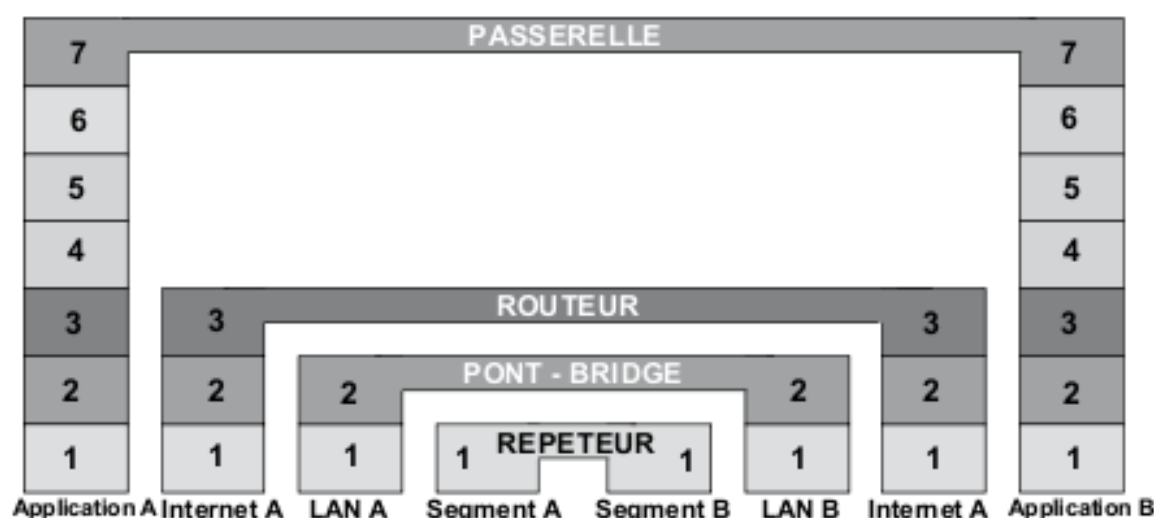


Figure 29.1 Les organes d'interconnexion selon l'ISO.

29.2 Les routeurs

29.2.1 Les routeurs et les passerelles inter-réseau

Un routeur relaie les paquets entre deux réseaux d'espace d'adressage homogène (IP/IP, ISO/ISO...). Lorsque l'espace d'adressage n'est pas homogène, par exemple interconnexion de réseaux IP *via* un réseau X.25, il est nécessaire de mettre en œuvre un mécanisme de conversion d'adresses (IP/ISO). Ce mécanisme n'est pas normalisé et chaque constructeur apporte sa solution. L'organe d'interconnexion n'est plus strictement un routeur, c'est une passerelle inter-réseau que le langage courant continue d'appeler routeur.

29.2.2 L'architecture d'un routeur

Un routeur met en relation un couple de ports d'accès (LAN ou WAN) identifiés par une adresse réseau. Le schéma de la figure 29.2 représente l'architecture simplifiée d'un routeur.

La configuration générale d'un routeur consiste en une succession de déclarations pour effectuer la mise en relation des différents modules et en un ensemble de paramètres décrivant leurs caractéristiques (protocole utilisé, MTU, taille fenêtre...). Un routeur n'est pas obligatoirement une machine spécifique, une station d'un réseau local peut participer à l'acheminement des données (Unix, Linux, Windows...).

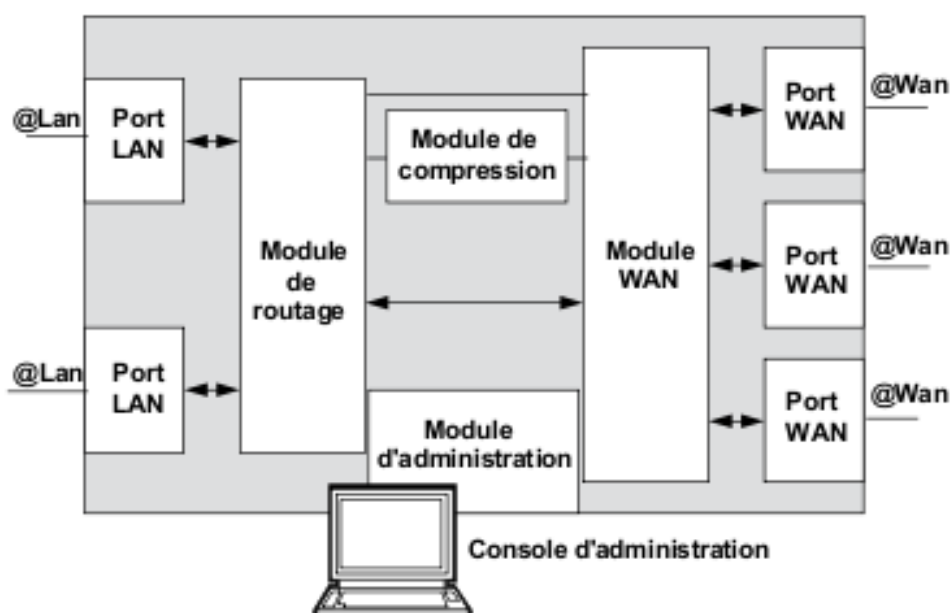


Figure 29.2 L'architecture simplifiée d'un routeur.

30

Les techniques de routage

30.1 Généralités

Rappelons-en le principe, les routeurs acheminent les paquets selon des informations contenues dans des tables dites tables de routage. Ils utilisent essentiellement deux modes de routage :

- ▶ le routage statique ou fixe, dans ce type de routage, les tables de routage sont introduites par l'administrateur de réseau à la configuration du réseau ;
- ▶ le routage au moindre coût, dans ce type de routage, les tables de routages indiquent pour chaque destination le chemin de coût le moins élevé. Périodiquement, des échanges d'information entre les routeurs permettent de maintenir ces tables à jour. Les algorithmes de vecteur distance (*Distance vector routing*) et ceux à état des liaisons (*Link state routing*) sont de cette nature.

La station qui a des données à transmettre connaît le routeur auquel elle est rattachée (routeur ou passerelle par défaut). Ce routeur doit ensuite déterminer le prochain nœud à atteindre pour trouver le destinataire. Ce choix est effectué par consultation de la table de routage en fonction d'une politique de routage. Un protocole de routage résout essentiellement trois problèmes :

- ▶ il découvre les autres routeurs du réseau,
- ▶ il construit les tables de routage,
- ▶ il maintient les tables de routage à jour.

Si on veut réaliser l'interconnexion de réseaux d'opérateurs différents, il est nécessaire, à l'interface entre les réseaux, de définir un protocole commun d'échange des informations de routage. Chaque opérateur

peut alors, en interne dans son réseau, utiliser le mode de routage qui lui convient. Aussi, outre l'aspect de limitation du trafic de gestion, le domaine global de routage (**Internet**) a été subdivisé en domaines de routage autonomes (**AS**, *Autonomous System*). Cette division conduit à distinguer deux familles de protocoles de routage (figure 30.1) :

- ▶ les protocoles de routage intra-domaine, pour le routage à l'intérieur d'un même domaine (**IGP**, *Interior Gateway Protocol*). Les paquets de service du protocole de routage identifient le domaine d'appartenance, tout paquet n'appartenant pas à ce domaine est ignoré. Cette technique limite la diffusion à l'intra-réseau ;
- ▶ les protocoles de routage inter-domaine (**EGP**, *Exterior Gateway Protocol*). Ces protocoles routent les paquets d'information dans l'inter-réseau. Ces protocoles doivent prendre en compte les accords commerciaux ou politiques entre les différents opérateurs des systèmes autonomes.

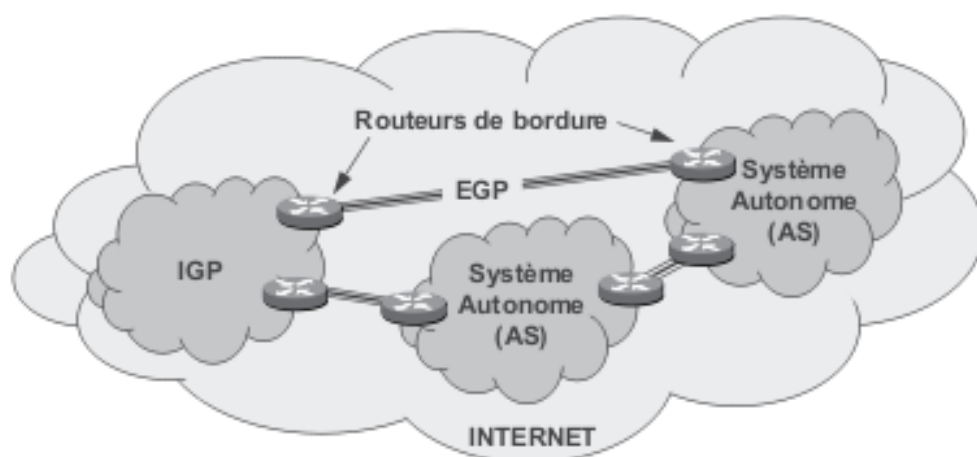


Figure 30.1 Le découpage d'Internet en domaines de routage.

30.2 Le routage dans le réseau IP

30.2.1 L'adressage d'interface

Supposons le réseau simplifié de la figure 30.2, peu importe le protocole de niveau 2 mis en œuvre sur le lien reliant les passerelles d'accès (routeurs IP). Comment la couche IP peut-elle déterminer l'interface de sortie par

rapport à une adresse IP destination alors que la couche IP ignore la technologie sous-jacente ?

En application du principe d'indépendance des couches, le point d'accès au réseau physique ne peut être connu de la couche IP que par une adresse IP. Les liaisons entre les différentes passerelles sont considérées, vu d'IP, comme constituant un réseau ; de ce fait chaque extrémité d'une liaison possède une adresse IP. Dans ces conditions, l'algorithme d'acheminement recherche sur quel réseau (de liaisons) est situé le saut suivant. Cette technique d'identification de l'interface d'accès au réseau physique est dite adressage d'interface ou **adressage de LS** (Liaison spécialisée). Cette méthode garantit l'indépendance des couches. En effet, le routage se réalise d'adresse IP destination à adresse IP d'interface (*Next hop*).

Les liens inter-routeurs forment ainsi le réseau logique IP. Une adresse IP est attribuée à chaque extrémité. Pour comprendre le mécanisme de routage, la figure 30.2 fournit un exemple de configuration d'un routeur (Routeur R1). Notons que la route à prendre est désignée par l'adresse distante du lien (*Next hop*), ce qui correspond à l'adresse du point à atteindre sur le réseau de liens.

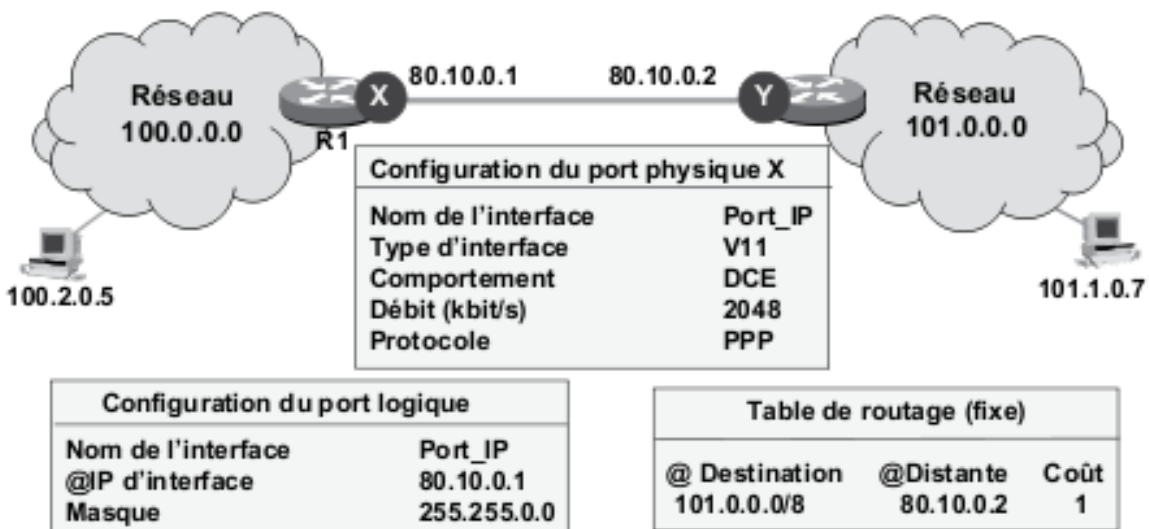


Figure 30.2 Exemple de configuration d'un routeur IP.

30.2.2 Le concept d'interface non numérotée

L'attribution d'adresses d'interface est consommatrice d'adresses, aussi le RFC 1812 a-t-il autorisé le routage sur interface dite non numérotée (*Unnumbered IP*). Un exemple de configuration simplifiée est donné par la figure 30.3. Cette approche viole la règle d'indépendance des couches. Aussi, le RFC 1812 précise que les deux passerelles connectées par une ligne point à point non numérotée ne sont pas à considérer comme deux routeurs mais comme deux demi-routeurs constituant un seul routeur virtuel (figure 30.3).



Figure 30.3 Le routage du RFC 1812.

L'acheminement de datagrammes IP entre deux réseaux locaux directement reliés par un routeur est illustré par la figure 30.4. En 1, la machine source doit envoyer un datagramme à la machine d'adresse IP @IPD. L'analyse du masque de sous-réseau lui montre que la machine cible n'est pas située sur le même réseau qu'elle, elle va donc faire appel aux services de la passerelle par défaut.

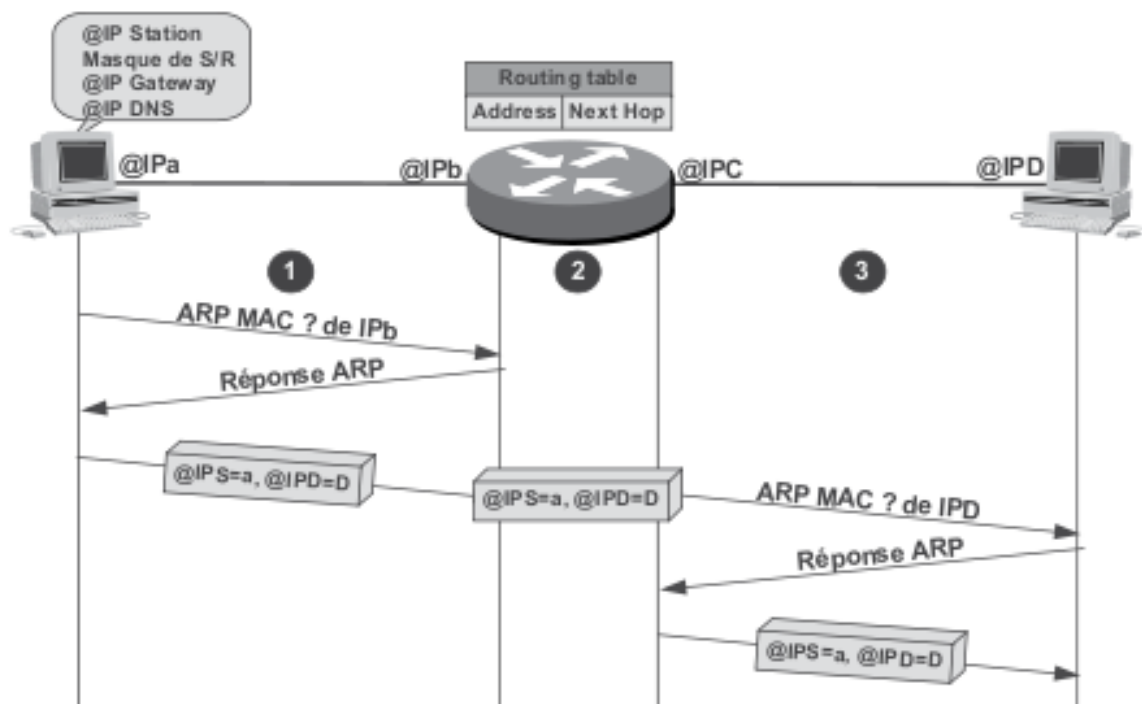


Figure 30.4 L'acheminement entre deux réseaux locaux.

Les données acquises lors de sa configuration lui fournissent l'adresse IP de la « passerelle » par défaut, mais pas son adresse MAC. Elle émet alors une requête ARP pour découvrir l'adresse MAC de cette « passerelle », puis lui transmet le message. En 2, le routeur recevant un datagramme dont l'adresse IP n'est pas la sienne, sait qu'il doit en réaliser l'acheminement. Il consulte alors sa table de routage qui lui indique l'interface de sortie. Cette interface est une interface de réseau local, pour acheminer le paquet il doit donc connaître l'adresse MAC du destinataire. Pour cela en 3, il émet une requête ARP sur ce segment local, puis achemine le datagramme vers son destinataire.

30.2.3 Le routage vecteur-distance, RIP

Issu des travaux de Bellman-Ford, le protocole **RIP** (*Routing Information Protocol*, RFC 1058), développé par l'Université de Californie (UCB, *University of California at Berkeley*) pour Unix BSD 4.2 (*Berkeley Software Distribution*) et utilisé initialement dans Arpanet, est en raison de sa simplicité, de sa solidité, de sa facilité de mise en œuvre et, ceci

malgré ses lacunes, le protocole de routage vecteur distance le plus utilisé.

RIP distingue deux types d'équipement : les actifs et les passifs. Les premiers diffusent périodiquement leurs routes vers les autres nœuds tandis que les seconds écoutent et mettent simplement leur table à jour en fonction des informations reçues. Un routeur fonctionnant en mode actif envoie toutes les 30 secondes un message de diffusion pour signaler qu'il connaît une route pour accéder à un réseau et il en indique le coût en nombre de sauts.

Dans les grands réseaux, la diffusion, toutes les 30 secondes, des tables de routage induit un trafic important et un temps de convergence important (stabilisation des tables) qui peut être de plusieurs minutes. Pour limiter ce temps, la visibilité d'un routeur est limitée à 15 sauts, une métrique de 16 représente une route non joignable. Si un routeur ne reçoit aucun message durant 180 secondes, la route silencieuse est déclarée inaccessible (coût = 16).

Le protocole **IGRP** (*Interior Gateway Routing Protocol*) de Cisco remédie à de nombreux inconvénients de RIP. Dans RIP, le coût est représenté par le nombre de sauts, une voie plus rapide (débit ou temps) est ignorée si son transit représente un nombre de sauts plus important. IGRP utilise une métrique, configurable par l'administrateur, qui permet de privilégier un lien.

30.2.4 Le routage à état des liens (OSPF)

■ Généralités

Contrairement au protocole à vecteur distance, le protocole à état des liens ne diffuse, sur le réseau, que les modifications qu'il a détectées dans la topologie du réseau : lien inaccessible, coût modifié... Chaque nœud entretient une base de données qui est le reflet total de la cartographie du réseau. Cette vision globale, par chaque routeur, du réseau permet d'éviter la formation de boucles. Le coût de la liaison (métrique) est configurable interface par interface, plusieurs métriques peuvent être utilisées simultanément (longueur de la file d'attente, débit, distance...). À partir de ces éléments, chaque routeur calcule la route de moindre coût selon l'algorithme de Dijkstra.

OSPF (*Open Shortest Path First*) est capable d'assurer un routage par type de service (champ TOS du datagramme IP), il peut aussi assurer l'équilibrage de charge entre plusieurs routes de même coût. Lorsque le réseau est important, la diffusion des messages et la détermination de la nouvelle table de routage pénalisent les performances globales du réseau. Aussi, OSPF a introduit la notion d'aire limitant ainsi l'espace de diffusion (nombre de routes à annoncer) et le volume de calcul à effectuer.

■ La notion d'aire de routage

Une aire ou zone (*Area*) correspond à une subdivision logique d'un réseau OSPF (figure 30.5). Il est important de ne pas confondre la notion d'aire d'OSPF avec celle de système autonome (**AS**, *Autonomus System*). Les protocoles de routage utilisés dans chacun des AS peuvent être différents, un protocole spécifique (**EGP**, *External Gateway Protocol*) gère l'échange d'information entre les différents AS alors que les différentes aires OSPF utilisent toutes le protocole OSPF.

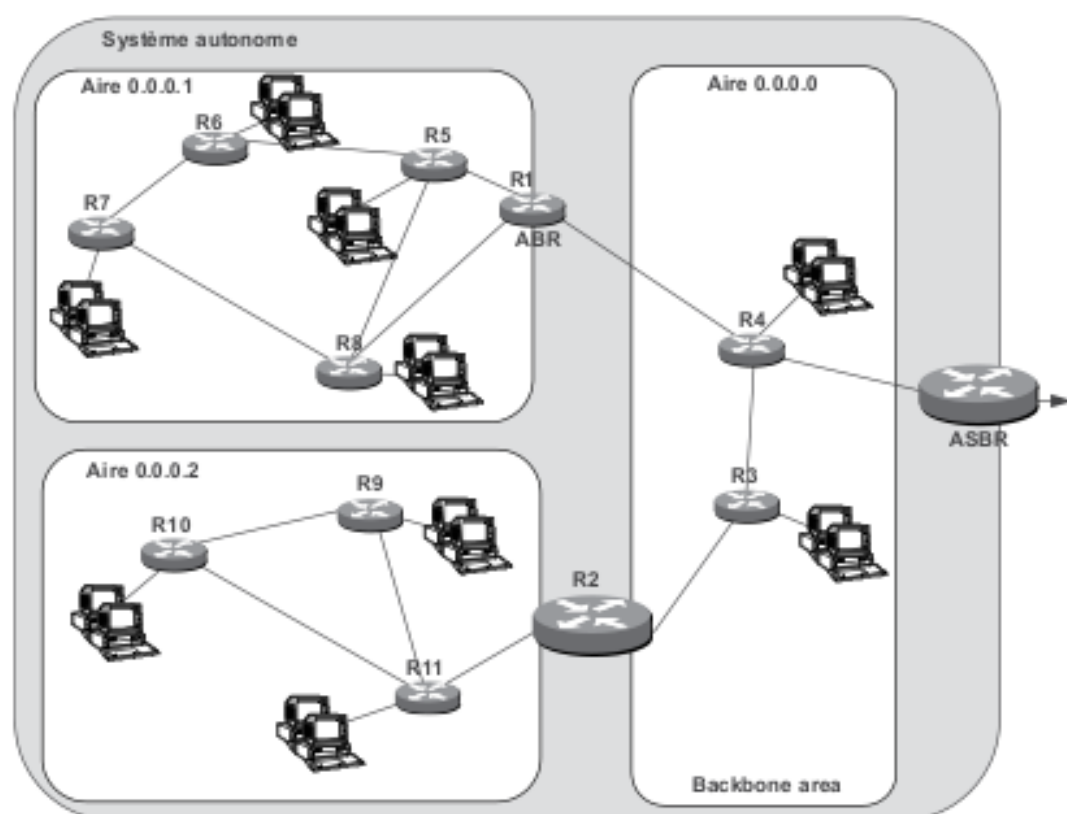


Figure 30.5 Les aires OSPF.

La hiérarchie introduite par OSPF est limitée à deux niveaux. L'environnement OSPF comprend les éléments suivants :

- ▶ une zone dite fédératrice (*Area backbone*) assure l'interconnexion de toutes les autres zones. Chaque zone est identifiée par un numéro de zone unique ;
- ▶ des routeurs de zone ou *Internal Router (IR)*, ces routeurs n'annoncent que les routes internes à leur zone ;
- ▶ des routeurs qui assurent la connexion au *backbone*, et qui annoncent les routes extérieures à la zone. Ainsi, sur la figure 30.5, il s'agit du routeur R1 (**ABR**, *Area Boundary Router*), le routeur R2 qui correspond à une interface locale entre la zone 2 et la zone 0 est considéré comme appartenant au *backbone* ;
- ▶ des routeurs frontières de système autonome (**ASBR**, *Autonomous System Boundary Router*). Ces routeurs assurent l'échange d'information avec les autres systèmes autonomes. Les routes extérieures au système autonome sont apprises par des protocoles autres qu'OSPF (routage statique, EGP, BGP...).

La réduction du nombre de routeurs par zone de diffusion limite le trafic de gestion mais les échanges entre routeurs sont encore nombreux. Pour limiter ce trafic, OSPF introduit la notion de routeur désigné (**DR**, *Designated Router*). C'est ce dernier qui assure la diffusion des messages vers les routeurs de la zone ce qui ne nécessite que N messages (1 message vers le DR et $N-1$ messages du DR vers les $N-1$ hôtes).

■ L'agrégation de routes

L'utilisation de zones offre un mécanisme puissant d'affectation des adresses IP. Si tous les réseaux ou sous-réseaux d'une zone ont des adresses IP contiguës, le routeur ne signale qu'une seule route aux autres routeurs. Cette propriété permet d'une part de minimiser le trafic d'annonce et, d'autre part, d'alléger les tables de routage. La figure 30.6 illustre ce principe.

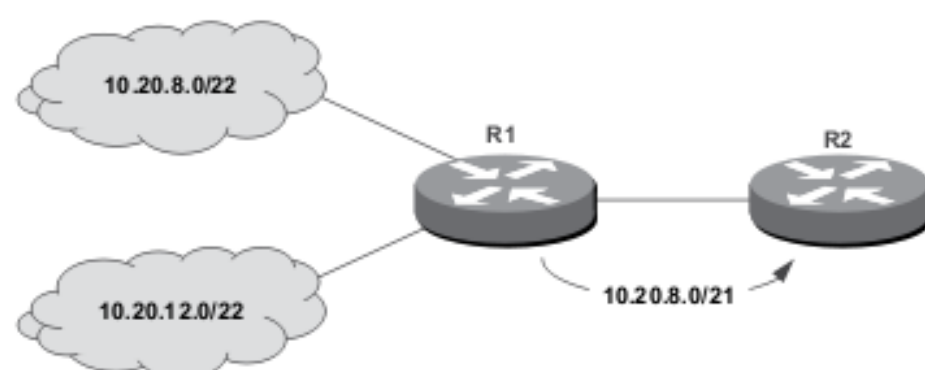


Figure 30.6 L'agrégation de routes dans OSPF.

En routage vecteur distance, les deux sous-réseaux de la figure 30.6 seraient annoncés individuellement et constitueraient 2 entrées dans les tables de routage de R1 et de R2. En OSPF, le routeur R1 considère que les réseaux ont en commun les 21 premiers bits, il n'annonce que cette seule information.

30.2.5 Routage et qualité de service

Lorsque les réseaux sont largement sur-dimensionnés, il n'est nullement besoin d'avoir recours à des mécanismes particuliers, les flux s'écoulent « librement ». Ce n'est que quand la ressource est insuffisante qu'il devient nécessaire de gérer la bande passante pour n'allouer à chaque flux que ce dont il a besoin en effectuant éventuellement des arbitrages, c'est la notion QoS (*Quality of Service*). La qualité de service ne crée pas de bande passante, elle la gère en prenant en compte les exigences de chaque flux et en les traitant de manière différenciée par rapport aux autres flux moins exigeants.

Conçu à une époque où les seuls flux applicatifs à acheminer étaient tous de même nature (texte), IP n'implémentait qu'un mécanisme simple de qualité de service défini selon les 3 bits du champ TOS. Le champ TOS, aujourd'hui obsolète, a été redéfini. Deux approches de la QoS sont actuellement définies : *Integrated Services* (**IntServ**, RFC 1633) peu mis en œuvre et *Differentiated Services* (**DiffServ**, RFC 2474).

■ Differentiated Services

Un réseau « DiffServ » met en œuvre un certain nombre de mécanismes symbolisés par la figure 30.7 :

- ▶ Il identifie chaque flux en fonction d'une règle de gestion de la bande passante (classification et marquage). La classification est établie en périphérie du réseau, ce qui « soulage » le réseau, mais elle est réalisée sans connaissance de l'état de réseau ;
- ▶ Il analyse le trafic pour gérer la bande passante et offrir à chaque flux un traitement différencié en fonction de sa classification ;
- ▶ Il supervise les différents éléments actifs du réseau pour prévenir la congestion en mettant préventivement en œuvre une politique d'élimination ;
- ▶ enfin, en cas de congestion, il applique une politique de distribution de la bande passante (gestion des files d'attente). Les mécanismes de QoS ne sont mis en œuvre qu'en cas de congestion, en l'absence de congestion, le traitement des files d'attente est du type FIFO (*First In, First Out* ; premier arrivé, premier servi).

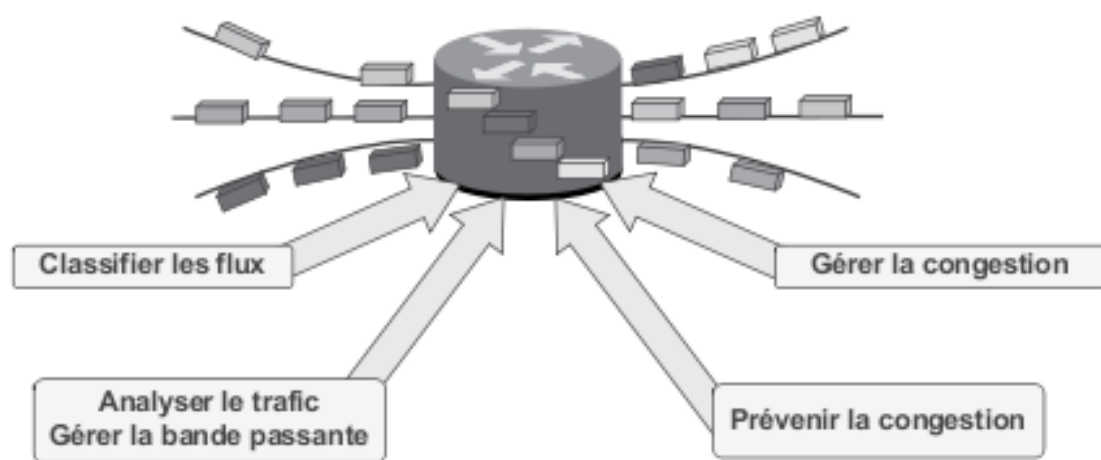


Figure 30.7 Ensemble des mécanismes d'un réseau DiffServ.

L'approche *DiffServ* (services différenciés) est plus conforme à l'approche IP. *DiffServ* implémente un mécanisme de partage de bande passante en introduisant la notion de politique d'acheminement en fonction d'une classe de service (RFC 2474, *Differentiated Service* ou *Diffserv*). Les flux ne

sont plus traités individuellement comme dans *IntServ*, mais sont affectés à une classe de service identifiée par un champ spécifique *Differentiated Services Field (DSF)*. Tous les flux d'une même classe sont traités de la même manière dans le réseau dit traitement par saut (**PHB**, *Per Hop forwarding Behaviour*). Le champ *DiffServ* remplace le champ TOS d'IPv4 et le champ Classe de Service d'IPv6. Voir la figure 13.1 du chapitre 13 qui rappelle la structure de chacun de ces champs.

Diffserv répartit le trafic en trois classes de service (**CoS**, *Class of Service*) :

- ▶ *Expedited Forwarding* ou *Premium Service* (RFC 2598), équivalent aux services CBR et VBR-rt d'ATM et défini spécifiquement pour les applications temps réel, minimise le temps de latence dans le réseau. Celui-ci prend en compte des contraintes fortes en matière de temps de traversée, de variation de celui-ci (gigue) et de pertes de données. L'utilisation de la classe *Premium* doit être limitée au sein du réseau, un abus d'utilisation peut fortement pénaliser les autres types de trafic (champ DSCP = 101DD0) ;
- ▶ *Assured Forwarding* ou *Olympic Service* (RFC 2597), équivalent des services ABR d'ATM, comporte 4 classes, elles-mêmes subdivisées en fonction d'une politique d'écartement (**RED**, *Random Early Drop*) en fonction de l'état du réseau (*Low drop*, *Medium drop* et *High drop*). À chaque classe sont affectées une priorité et une garantie de bande passante ;
- ▶ *Best Effort*, équivalent du service UBR d'ATM, correspond au trafic traditionnel sur IP sans qualité de service (champ DSCP = 000000).

La classification des flux est réalisée à la périphérie du réseau (*Classifier*). Les propriétés du flux sont ensuite analysées (*Metering* ou dimensionnement) en fonction d'un contrat de service préétabli (**SLA**, *Service Level Agreement*). Les datagrammes sont alors marqués (*Marker*) par positionnement du champ **DS Field**. Le trafic différencié ou « colorisé » (*Multiflow Classifier*) est analysé, certains paquets peuvent être retardés (mise en forme du trafic ou *shaper*), voire éliminés pour prévenir un éventuel état de congestion (*Dropper*). Les paquets sont ensuite affectés à une file d'attente spécifique avant d'être transmis sur le réseau (*Forwarding*). La figure 30.8 illustre les mécanismes que nous venons de décrire.

L'architecture *DiffServ* est bien adaptée aux grands réseaux. En effet, les classes de services étant attribuées en périphérie du réseau *DiffServ*, elles ne génèrent ni trafic de gestion, ni surcharge CPU. Cependant, si *DiffServ* permet de hiérarchiser les flux, il ne dispose pas de mécanisme d'information d'état du réseau. De ce fait, les routeurs de bordure ne sont pas en mesure d'anticiper ni de réagir à un état de précongestion ou de congestion.

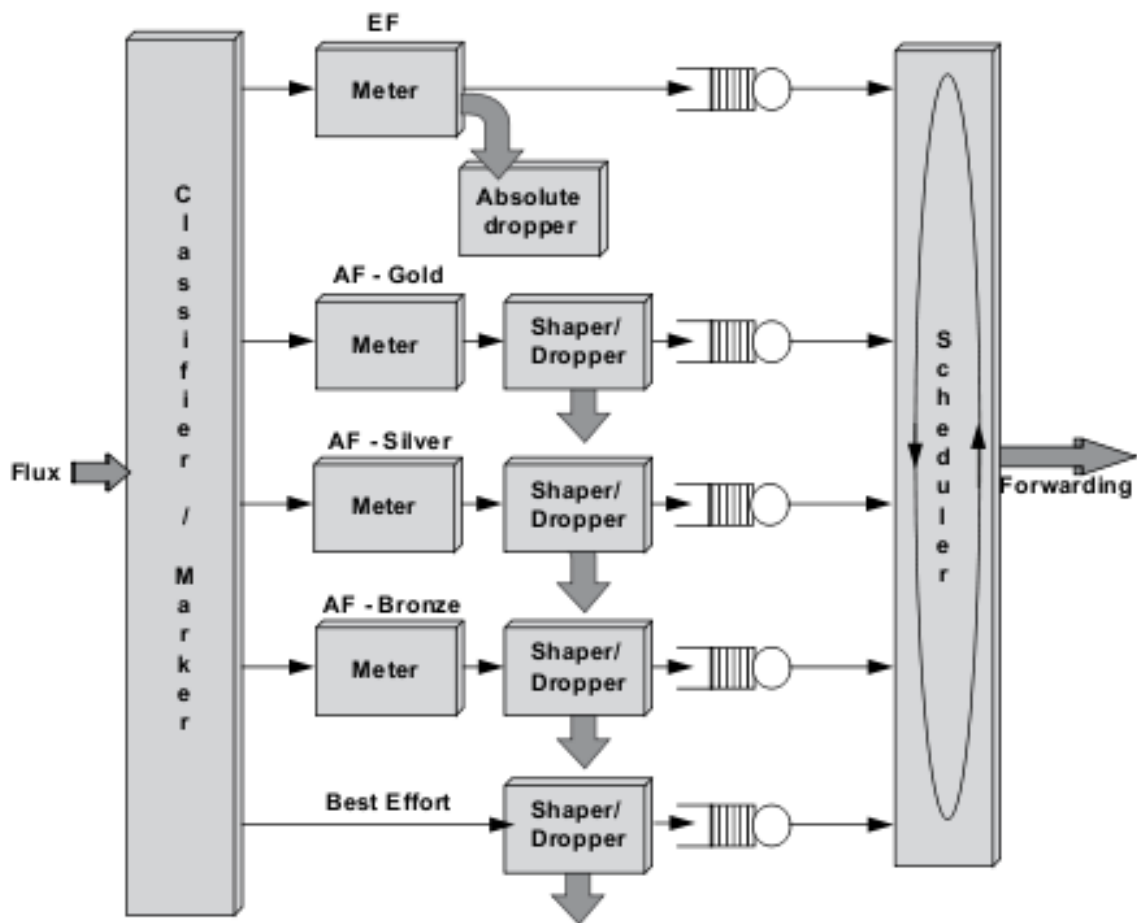


Figure 30.8 Traitement des flux sous DiffServ.

31

Le routage *multicast*

31.1 Introduction au *multicast*

L'IP *multicast* est un mécanisme qui permet de diffuser des datagrammes IP vers un ou plusieurs récepteurs sans que ces datagrammes soient adressés individuellement à chaque hôte. Les nœuds de destination sont identifiés par une adresse de groupe (adresse *multicast*). Cette technique évite l'envoi de N datagrammes *unicast* aux N clients d'une même source d'information.

Dans l'exemple de la figure 31.1, un serveur d'applications vidéo (serveur *multicast*) est raccordé à un réseau IP *multicast*.

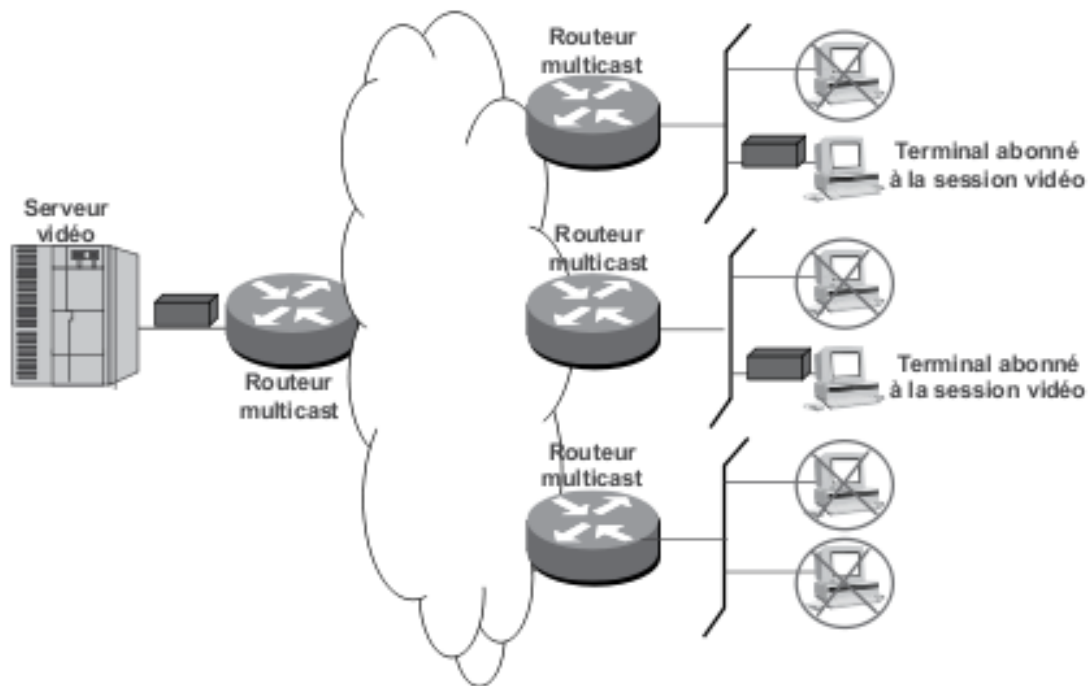


Figure 31.1 Principe d'un réseau multicast.

Pour recevoir les émissions, les machines distantes doivent, au préalable, s'abonner au serveur vidéo. Le mécanisme d'abonnement est spécifique au logiciel serveur. Les machines abonnées sont membres du groupe *multicast*. Une machine peut être membre de plusieurs groupes. Les informations ne sont diffusées sur les brins locaux que si au moins une machine locale s'est abonnée à un service *multicast*. Dans notre exemple, un seul exemplaire du datagramme est diffusé sur le réseau *multicast*, chaque routeur ayant un client abonné rediffuse le datagramme sur le seul brin local auquel au moins un client est raccordé. Sur le troisième réseau où aucune machine n'a rejoint un groupe de *multicast*, les informations vidéo ne sont pas diffusées.

31.2 Le protocole local IGMP (RFC 2236)

Le protocole **IGMP** (*Internet Group Management Protocol*) est le protocole d'apprentissage utilisé par les routeurs *multicast* pour découvrir l'existence, dans les sous-réseaux auxquels ils sont raccordés, de membres d'un groupe *multicast* (figure 31.2).



Figure 31.2 Principe du protocole IGMP.

31.3 Les protocoles de routage *multicast*

31.3.1 Généralités

IGMP assure la diffusion des datagrammes *multicast* à l'ensemble des stations d'un sous-réseau local. Cependant, dans le réseau (Internet ou autre), le routage des datagrammes *multicast* doit être assuré de manière

à économiser la bande passante. Plusieurs mécanismes peuvent être mis en œuvre :

- ▶ le *Flooding* ou inondation, très facile à mettre en œuvre puisque le routeur n'a aucune table d'acheminement à entretenir, mais cette méthode génère un trafic d'autant plus important que le réseau est grand ;
- ▶ le *Spanning Tree* permet de construire un chemin unique (arbre) entre une source (racine) et une destination (feuille) ;
- ▶ le *Reverse Path Broadcasting (RPB)* aussi appelé *Reverse Path Forwarding* est une amélioration des performances du *Spanning Tree*. RPB construit un arbre par groupe de *multicast*. Le principe est relativement simple quand un routeur reçoit un paquet sur une interface, il examine si l'interface d'arrivée est bien sur le chemin le plus court pour rejoindre la source (celle qui serait utilisée pour envoyer un datagramme *unicast* à la source). Si c'est le cas, le paquet est diffusé sur toutes les autres interfaces du routeur (inondation), sinon le paquet est détruit, de proche en proche l'arbre est construit.

31.3.2 DVMRP, protocole de routage multicast

DVMRP (*Distance Vector Multicast Routing Protocol*, RFC 1075) est un protocole vecteur distance destiné à assurer le routage *multicast*. Une variante de l'algorithme *Reverse Path Broadcasting* est utilisée pour construire les tables de routage *multicast*. DVMRP est le complément réseau du protocole local IGMP (figure 31.3).



Figure 31.3 La complémentarité des protocoles IGMP et DVMRP.

L'arbre *multicast* établit des relations en « amont » et en « aval » entre les différents routeurs du réseau. L'arbre de diffusion est maintenu à jour

selon une méthode dite élagage et greffe. Lorsqu'un routeur *multicast* n'a plus d'abonné pour un groupe (hôte final ou routeur aval), il émet un message dit *prune packet* pour informer en amont de l'inutilité de maintenir un lien de diffusion *multicast* pour ce groupe (élagage de la branche). À l'inverse à la réception d'une demande d'adhésion, il émet un message dit *graft packet* pour reconstruire une route vers la source (greffe). Ce fonctionnement est illustré figure 31.4.

À l'instar de RIP, le protocole DVMRP échange ses tables de routage avec ses voisins DVMRP (message *route report* toutes les 60 s à l'adresse de diffusion 224.0.0.4). Les voisins DVMRP sont découverts par l'émission, toutes les 10 s, de messages *probe message*.

Issu des travaux du groupe IDRM (*Inter-Domain Multicast Routing*) de l'IETF, le protocole **PIM** (*Protocol Independant Multicast*) est essentiellement destiné à être mis en œuvre sur les grands réseaux comme Internet qui se caractérisent par des serveurs et des clients très dispersés. PIM se décline en deux versions, PIM DM (*Dense Mode*) et le PIM SM (*Sparse Mode*, RFC 2362).

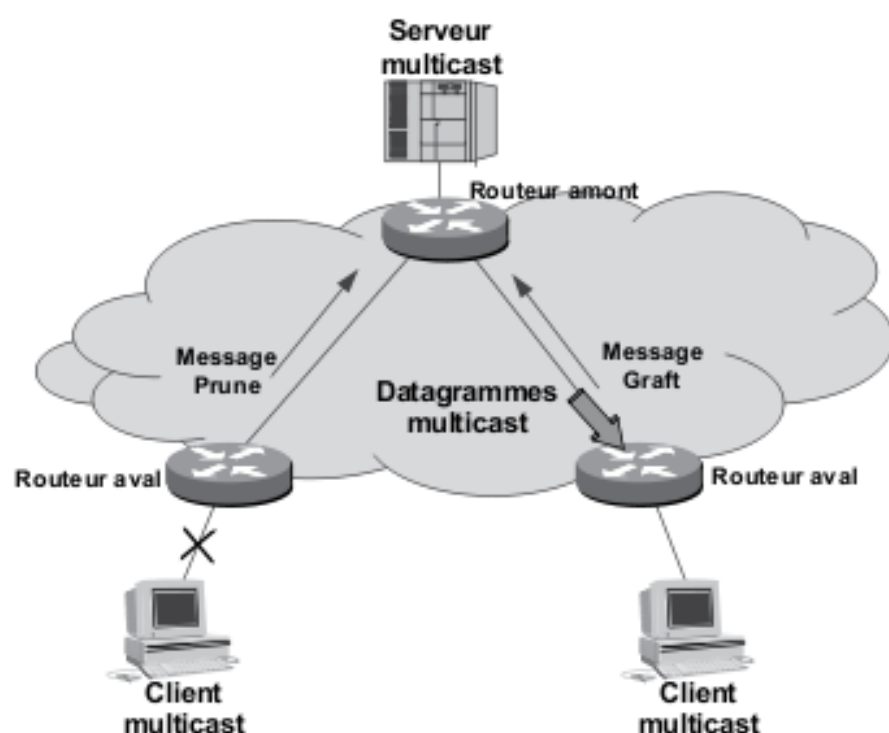


Figure 31.4 La diffusion de messages dans un réseau DVMRP.

Le PIM SM intègre son propre algorithme de routage *unicast* (indépendant du protocole de routage *unicast* utilisé dans le réseau). Les routeurs d'un sous-réseau en charge de la gestion du *multicast* doivent explicitement rejoindre un groupe *multicast* en s'abonnant à un point de rendez-vous pour ce groupe (message « *Join* »). À un groupe *multicast* correspond un seul point de rendez-vous.

31.4 Internet et le *multicast*

La composante d'Internet qui assure la diffusion de messages en *multicast* sur le réseau est désignée sous le nom de **Mbone** (*Multicast Backbone*). Il s'agit d'un réseau virtuel reliant les différents routeurs *multicast* (mrouteurs) par des tunnels (tunnels *multicast*), mais le service rendu est du type datagramme.

10

**La téléphonie
sur IP**

32

Principes généraux de la téléphonie

32.1 Introduction

Historiquement, le transport de la voix est à l'origine des premiers réseaux de transmission. Utilisant le principe de la commutation de circuits, le réseau téléphonique public commuté (**RTPC**, ou simplement **RTC**, ou encore **PSTN** pour *Public Switched Telecommunication Network*) met en relation deux abonnés à travers une liaison dédiée pendant tout l'échange (figure 32.1).

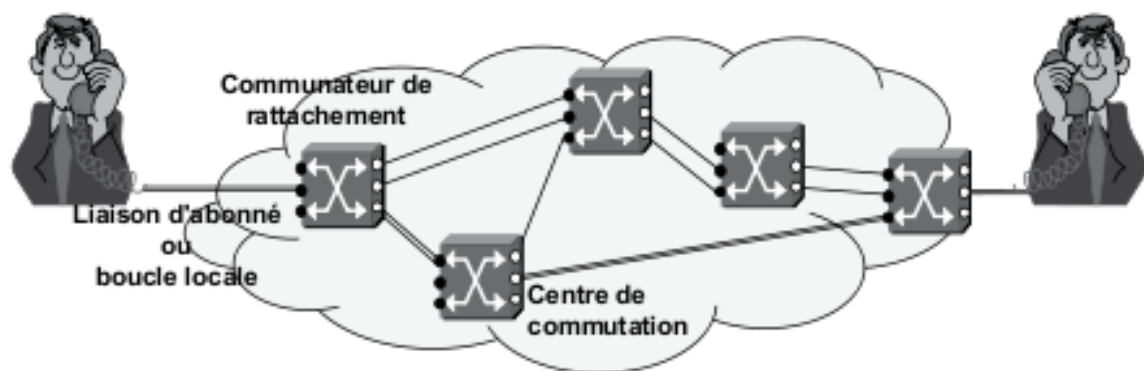


Figure 32.1 Principe du réseau téléphonique commuté.

La commutation de circuits ou commutation spatiale consiste à juxtaposer bout à bout des voies physiques de communication, la liaison est maintenue durant tout l'échange. La numérisation de la voix a permis de migrer le réseau d'une commutation spatiale vers une commutation d'intervalles de temps (IT) ou commutation temporelle (multiplexage).

32.2 De l'analogique à la ToIP

32.2.1 Du réseau analogique au réseau numérique

Le terminal restant analogique, la numérisation du réseau nécessite une conversion analogique/numérique en entrée du réseau et numérique/analogique en sortie. Cette conversion est réalisée par un convertisseur appelé **codec** (codeur/décodeur). Un usager qui désire pouvoir établir n communications téléphoniques simultanées doit être raccordé par n lignes (lignes groupées, les lignes groupées sont vues, pour le réseau, sous un même numéro). La numérisation autorise très simplement le multiplexage, d'où l'idée de réaliser des liaisons numériques de bout en bout, une seule ligne physique peut alors acheminer plusieurs communications téléphoniques (figure 32.2).

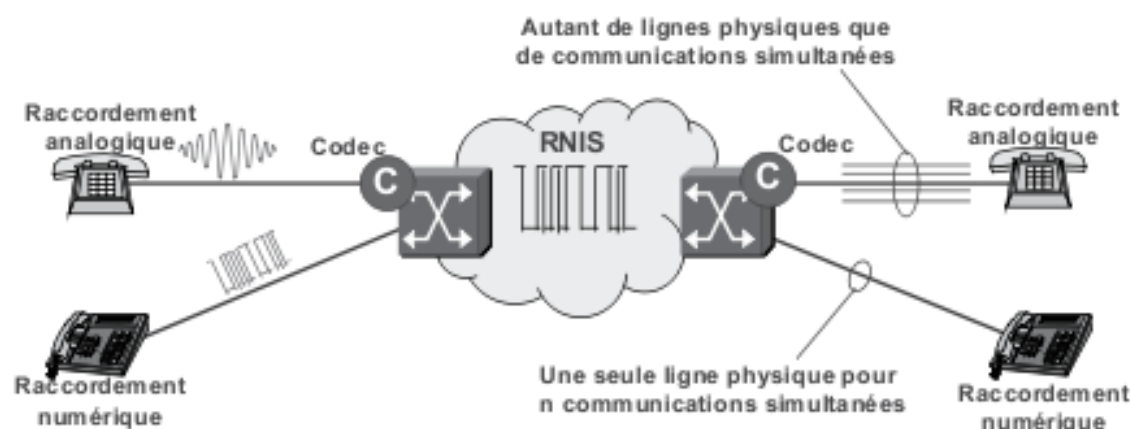


Figure 32.2 De l'analogique au numérique.

En réservant un IT (Intervalle de temps) à la signalisation (débit de 64 kbit/s), on peut l'acheminer, *via* un protocole de haut niveau, en mode message (protocole D). L'indépendance entre le flux média et le protocole de signalisation autorise de nombreux services nouveaux, c'est le **RNIS** (Réseau numérique à intégration de service ou **ISDN**, *Integrated Service Digital Network*).

32.2.2 La numérisation de la voix

■ Principe

Numériser une grandeur analogique consiste à transformer la suite continue de valeurs en une suite finie de valeurs discrètes. La figure 32.3 représente les différentes étapes de la numérisation du signal. À intervalle régulier (période d'échantillonnage), on prélève une fraction du signal (échantillon). Puis, on fait correspondre à l'amplitude de chaque échantillon une valeur discrète (quantification dite **quantification scalaire**) ; cette information est ensuite transformée en valeur binaire (codification).

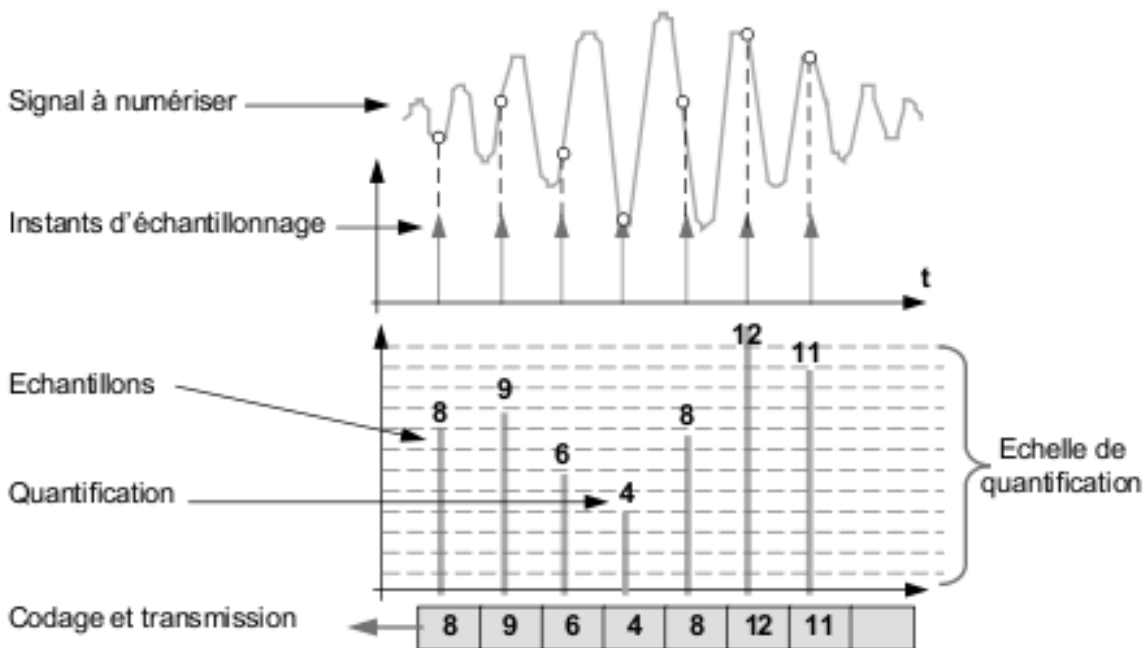


Figure 32.3 La numérisation d'un signal analogique.

Pour reproduire correctement le signal à l'arrivée, le récepteur doit disposer d'un minimum d'échantillons. Claude Shannon¹ a montré qu'il existait une relation étroite entre la fréquence maximale d'un signal et le nombre d'échantillons à prélever (relation de Shannon) :

$$F_{\text{échantillon}} \geq 2 \cdot F_{\text{max du signal}}$$

1 Claude Shannon (1916-2001) est l'un des initiateurs de la théorie de l'information.

Indépendamment du nombre d'échantillons, une reproduction correcte du signal nécessite que le système de transmission respecte l'espacement temporel entre échantillons. Les flux de données qui nécessitent une récurrence temporelle stricte (gigue nulle) sont dits isochrones et par abus de langage, flux ou données temps réel.

■ Application à la voix

Un canal téléphonique utilise une plage de fréquences ou largeur de bande s'étendant de 300 Hz à 3 400 Hz. En fixant à 4 000 Hz la fréquence de coupure du filtre d'entrée, la fréquence d'échantillonnage minimale est de :

$$F_e \geq 2 \cdot F_{\max} = 2 \cdot 4\,000 = 8\,000 \text{ Hz}$$

Soit 8 000 échantillons par seconde, ce qui correspond à prélever un échantillon toutes les 125 μ s (1/8 000). Pour une restitution correcte (dynamique¹ et rapport signal sur bruit), la voix devrait être quantifiée sur 12 bits (4 096 niveaux). L'utilisation d'une loi de quantification logarithmique permet de ramener la représentation numérique de la voix à 8 bits soit un débit de 64 kbit/s. Cette opération dite de compression est différente en Europe (loi A) et en Amérique du Nord (loi μ).

■ Le codage de la voix

Les procédés de codage et de compression de la voix déclinent trois techniques :

- ▶ Le codage **MIC** (Modulation par Impulsion et Codage ou **PCM**, *Pulse Code Modulation*) qui utilise une quantification logarithmique ;
- ▶ Les codages différentiels, codant non plus l'échantillon mais son écart par rapport à l'échantillon précédent comme l'**ADPCM** (*Adaptive Differential Pulse Code Modulation*) ;
- ▶ Des techniques plus élaborées reconstituent la voix par synthèse (**CELP**, *Code Excited Linear Prediction*).

¹ La dynamique exprime le rapport entre les puissances maximale et minimale du signal.

□ Appréciation de la qualité de la voix

Les algorithmes de compression sont destructifs. En l'absence d'instrument de mesure pour évaluer la qualité de la restitution sonore, celle-ci est directement appréciée par des observateurs humains. L'appréciation donnée est donc très subjective. La qualité est exprimée en **MOS** (*Mean Opinion Score*) sur une échelle de notes de 1 à 5 (tableau 32.1).

Tableau 32.1 L'expression du MOS selon l'ITU.

MOS	Appréciation de l'écoute	Fatigue ou effort à l'écoute
5	Excellent, pas de déformation perceptible	Aucun effort
4	Bon, dégradation à peine perceptible	Pas d'effort ni de fatigue importants
3	Passable, dégradation perceptible, mais on reconnaît son interlocuteur	Effort modéré nécessaire
2	Médiocre, dégradation considérable, la voix a une tonalité de voix synthétique	Effort et fatigue importants
1	Mauvais, quelques problèmes d'intelligibilité	Inintelligible

Le tableau 32.2 compare les différents algorithmes de compression en fonction du débit qu'ils autorisent et de la qualité de restitution de la parole. La norme G.711 est utilisée dans la téléphonie fixe traditionnelle. La norme G.729, mise en œuvre dans la voix sur IP, modélise la voix humaine par l'utilisation de filtres.

Tableau 32.2 Synthèse des principaux algorithmes de compression du son.

Codec	Technique	Débit	Délai	MOS	Commentaire
G.711	PCM	64 kbit/s	0,125 ms	4,2	Compression logarithmique. Téléphonie et H.320.
G.721	ADPCM	32 kbit/s	0,125 ms	4,4	Échantillon codé sur 4 bits.
G.722	ADPCM	48, 56, 64 kbit/s	< 3 ms		Spectre sonore de 50 à 7 000 Hz. Utilisation visioconférence (H.320).

Codec	Technique	Débit	Délai	MOS	Commentaire
G.722-1	ADPCM	16, 24, 32 kbit/s			Version 16 kbit/s supportée par Windows Messenger.
G.723	ADPCM	24 kbit/s	0,125 ms	3,9	
G.723-1	ACELP, MP-MLQ	5,2 et 6,4 kbit/s	30 ms	3,9	H.323 et H.324M (UMTS).
G.726	ADPCM	16, 24, 32, 40 kbit/s	< 3 ms	4,5 à 2,0	Décomposition du spectre en sous-bande. Remplace G.721.
G.727	ADPCM	16, 24, 32, 40 kbit/s		3,9	Complète G.726. Changement du nombre de bits (débit) en cours de conversation.
G.728	LD-CELP	16 kbit/s	3 ms	4,1	Trame courte de 2,5 ms
G.729	CS-ACELP	8 kbit/s	15 ms	3,9	Concurrent de G.723.1, recommandé pour voix sur Frame Relay et H.323.
G.729 a	CS-ACELP	8 kbit/s	15 ms	3,7	G.729 + récupération des silences.
GSM	RPE-LTP	13, 13,2 kbit/s	20 ms	3,9	

32.2.3 Du mode circuit au mode paquet

Transporter la voix sur un réseau en mode paquet (figure 32.4) nécessite de modéliser le flux voix à l'identique d'un flux de données en transformant le flux d'information continu (échantillons de voix) en un flux périodique (paquets). Si on peut garantir un transfert respectant les contraintes temporelles du transfert isochrone, il est alors possible d'utiliser un réseau à commutation de paquets pour transmettre la voix.

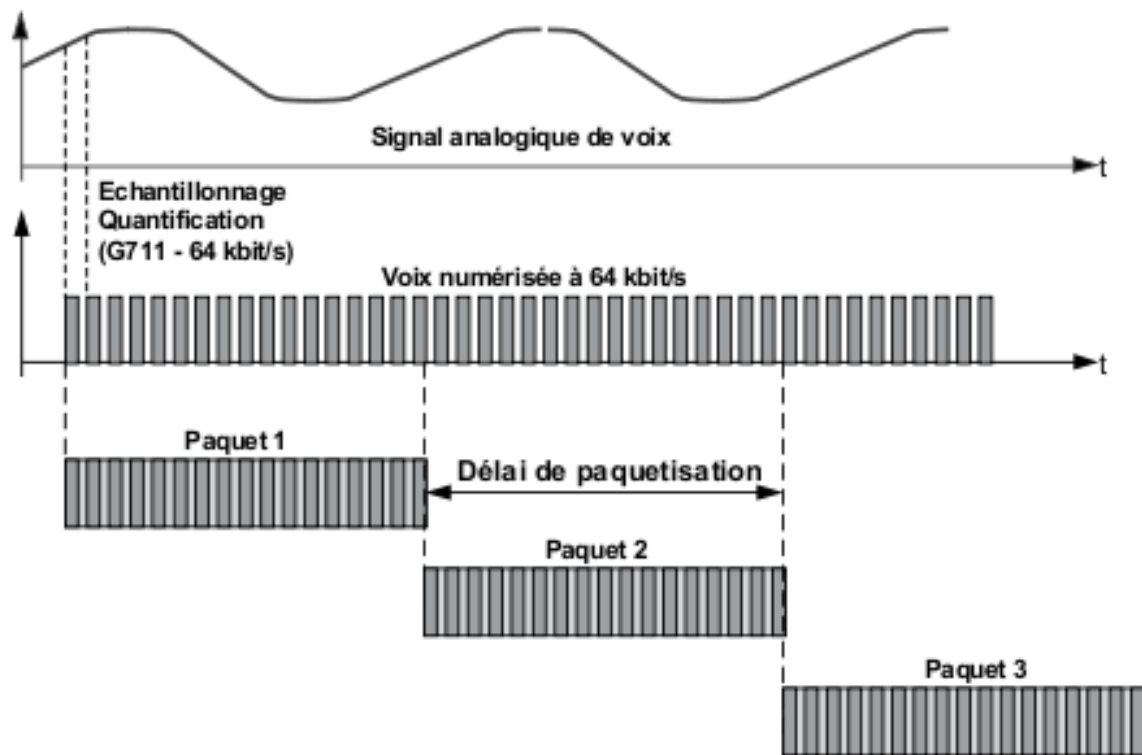


Figure 32.4 Du mode échantillon au mode paquet.

Pour garantir le transport de la voix, les systèmes d'interconnexion de réseaux voix/données doivent assurer :

- ▶ un débit minimal garanti, les flux voix ne devant jamais occuper plus de 40/50 % de la bande passante du support ;
- ▶ un transfert dans un délai de bout en bout aussi réduit que possible par priorisation des flux voix (< 150 ms, avec une tolérance jusqu'à 200 ms) ;
- ▶ la maîtrise de la gigue dans le réseau (< 40 ms, avec une tolérance jusqu'à 70 ms) et sa correction dans la passerelle destination.

32.2.4 Du numérique à la ToIP

La numérisation de la voix ayant permis la banalisation des flux, le transport mutualisé de la voix et de la donnée fut réalisé sur une même infrastructure de transport (multiplexage temporel ou TDM). Cependant le

service « données » et voix restaient assurés par des réseaux différents : le réseau téléphonique pour la voix ou RTC (Réseau téléphonique commuté ou encore PSTN, *Public Switched Telephon Network*) et des réseaux en mode paquet pour la donnée, historiquement avec le protocole X.25 (figure 32.5). Le paquetisation de la voix a, non seulement permis le transport en mode paquet, mais avec la voix sur IP la fusion totale des infrastructures réseaux du WAN au LAN.

Solutions d'avant hier (réseaux publics, réseaux privés)



Solutions d'hier (publiques ou privées)



Solution d'aujourd'hui : la voix sur IP



Figure 32.5 Évolution des réseaux.

La voix sur IP apparaît comme une nouvelle approche de la téléphonie, en fait les premières études du transport de la voix sur un réseau en mode paquet datent de 1972 (Bob Kahn). Les premières spécifications apparaissent avec le protocole NSC (*Network Secure Communications*) qui décrit comment transporter la voix sur ARPANET (1974).

En 1995, la société israélienne Vocaltec lance le premier logiciel de communication vocale de PC à PC : Internet phone. Ce mode de communication fut popularisé par Skype apparu dans le milieu des années 1990.

32.3 Notions d'autocommutateurs privés

32.3.1 Architecture

Un autocommutateur privé de téléphonie (figure 32.6), **PABX** (*Private Automatic Branch eXchange*) est l'interface (passerelle voix et passerelle de signalisation) entre le service de téléphonie de l'entreprise et un réseau téléphonique (public ou privé). Il est chargé de mettre en relation, à la demande de l'un des correspondants, des terminaux téléphoniques (commutation de circuits), que les deux correspondants soient derrière le PABX (communication interne à l'entreprise) ou que l'un d'eux soit sur un autre réseau téléphonique public (RTC ou RNIS) ou privé.

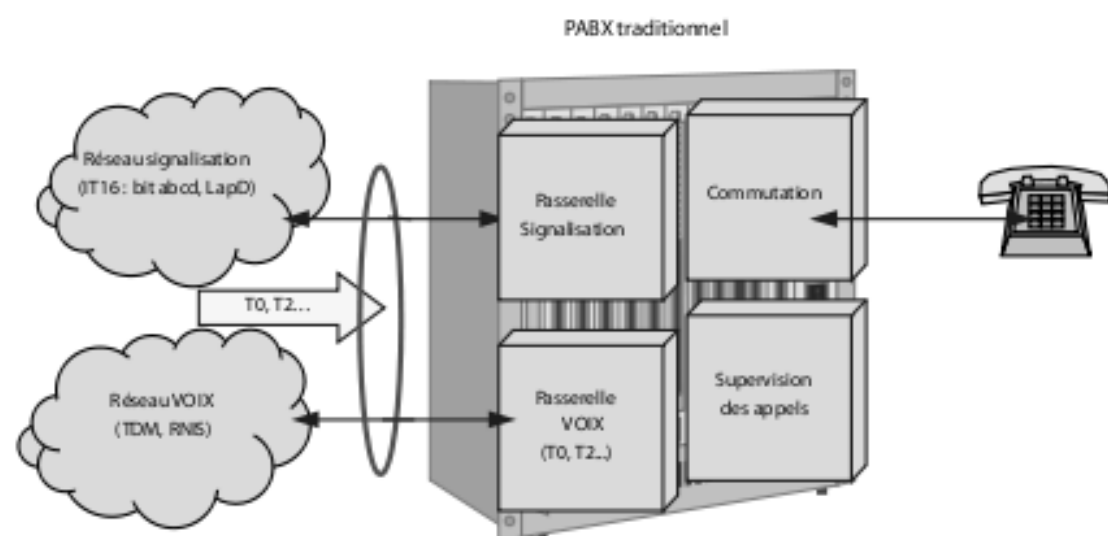


Figure 32.6 Architecture simplifiée d'un PABX.

La variété des services offerts par les PABX a conduit à les intégrer de plus en plus au système d'information de l'entreprise, l'aboutissement de cette évolution se concrétise par l'introduction de la voix sur IP.

32.3.2 Les facilités téléphoniques

Hors évidemment la réalisation d'une communication vocale, un système de téléphonie offre un ensemble de services spécifiques à la téléphonie désigné sous le terme de facilités. Le tableau 32.3 liste ces principales facilités, leur mise en œuvre et appellation diffèrent selon les constructeurs mais les finalités restent les mêmes.

Tableau 32.3 Exemples de facilités téléphoniques.

Facilité	Description
Annuaire	
Annuaire global	Consultation de l'annuaire global de l'entreprise. Eventuellement établissement d'appel par « click to call ».
Annuaire personnel	L'annuaire personnel peut comporter un extrait de l'annuaire global et de numéros propre à l'utilisateur du poste.
Poste multi-annuaire	Un poste multi-annuaire est un poste doté de numéros différents et offrant généralement des services différents selon le numéro appelé depuis l'extérieur (filtre, sonnerie) ou du numéro à partir duquel on réalise l'appel (discrimination accès externe ou interne, taxation...).
Appel par nom	Quand le poste est doté d'un clavier alpha numérique
Numérotation abrégée	Mise en correspondance d'un numéro court à signification interne avec un numéro long (intranet ou numéro normal externe). La numérotation abrégée peut être globale à l'entreprise (annuaire d'entreprise) ou individualisée par le titulaire du poste.
Renvoi	
Renvoi inconditionnel	Le poste appelé réachemine l'appel vers un autre poste interne (intranet), externe (renvoi dit en aboutement), la messagerie vocale ou sur le standard. La communication est renvoyée quel que soit l'état du poste appelé. Le renvoi inconditionnel peut être fixe (paramétrage du gestionnaire d'appels) ou variable (programmé par l'utilisateur).

Facilité	Description
Renvoi conditionnel	Le renvoi conditionnel, le renvoi n'est réalisé que si certaines conditions sont remplies : occupation, non réponse (nombre de sonneries), identification de l'appelant...
Transfert	Au cours de la communication, l'appelé peut rediriger la communication vers un autre abonné.
Identification	
N° appelant	Appel entrant, possibilité d'afficher sur le poste appelé le numéro d'appelant.
Nom appelant	Appel entrant, en association avec l'annuaire, mise en correspondance du N° appelant et du nom de celui-ci. Appel sortant, mise en correspondance du numéro appelé et du nom de l'appelé.
Restriction d'affichage	Appel sortant, interdiction appel par appel ou pour tous les appels de l'affichage de l'appelant. Généralement en entreprise cette information est alors remplacée par le N° du standard.
Supervision	Et service en dépendant.
Supervision de poste	Affichage sur le poste utilisateur de l'état d'un poste distant (libre, en communication, en sonnerie...)
Filtrage patron secrétaire	Secrétaire : poste filtrant les appels. Patron : poste filtré. Le poste filtrant supervise l'état du poste filtré et intercepte toutes ses communications. L'activation/désactivation de la fonctionnalité pouvant être réalisée aussi bien par le filtreur que le filtré. Le filtre peut être mis en œuvre pour tous les appels, ou différencié selon les types d'appel. Le filtrage peut être patron/secrétaire, plusieurs patrons/une secrétaire, un seul patron/plusieurs secrétaires...
Gestion des appels	
Sonnerie différenciée	En fonction de l'origine de l'appel (interne, externe), de la ligne appelée (poste multi-annuaire)
Repli de classe	Modification des droits du poste (appel interne, externe, renvoi.) sur demande du titulaire du poste ou automatiquement sur créneaux horaires.
Rappel automatique	Appel de l'appelé suite à un appel échoué : – Sur occupation du poste appelé ; – Sur saturation du faisceau de sortie.

Facilité	Description
Appel en attente	Indication du ou des appels en attente.
Message en attente	Indication d'un message vocal non lu.
Taxation	La taxation est un service de contrôle et de facturation en interne des appels. À chaque appel émis ou reçu le système émet une liste d'informations (ticket de taxation) contenant (liste non exhaustive) : - Le numéro appelé ; - Numéro appelant ; - Horodatage début communication (Heure/Date) ; - Horodatage fin communication ; - Durée de l'appel ; - Passerelle ayant établi l'appel ; - Nombre de sonneries avant décrochage ; - Transfert ; - Renvoi (numéro et type) ; - Type de communication...
Gestion centralisée	À chaque communication les « tickets » sont remontés et gérés sur un seul site de l'intranet
Gestion décentralisée	Chaque site gère ses communications
Division en sociétés	Répartition des postes en « services » pour une gestion différenciée.
Consultation	Possibilité par l'utilisateur de consulter ses tickets de taxation.
Masquage	Obligation par la CNIL du masquage des 4 derniers chiffres, obligation pouvant être levée par le titulaire du numéro.
Conférence	
Initialisation	Possibilité lors d'une communication d'insérer un (ou plusieurs) autre(s) correspondant(s).
Programmation	Programmation d'une conférence : intervenants, horaires.
Liste des participants	Possibilité de visualiser la liste des intervenants.
Retrait d'un correspondant	Volontaire par le participant ou autoritaire par un tiers participant.
Divers	
Musique d'attente	Possibilité de personnalisation.
Mise en attente	Interruption de la communication, puis reprise ;

Facilité	Description
Liste des appels en absence	Mémorisation et consultations des appels durant l'absence.
Interception dans le groupe	Appel reçu sur tous les téléphones du groupe (société) prise de l'appel par un téléphone appartenant au groupe, arrêt de sonnerie sur les autres postes.
...	

33

La téléphonie sur IP

33.1 Généralités

33.1.1 ToIP ou VoIP ?

La voix sur IP (VoIP) consiste à transporter la voix sous forme de paquets sur un réseau IP, alors que la téléphonie sur IP (ToIP) est un service téléphonique complet appuyé sur des équipements IP (téléphones...).

33.1.2 La ToIP

La téléphonie sur IP consiste simplement au remplacement du PABX d'établissement par un système de téléphonie reposant sur le protocole IP. Dans la figure 33.1, le système de téléphonie (établissement des appels...)

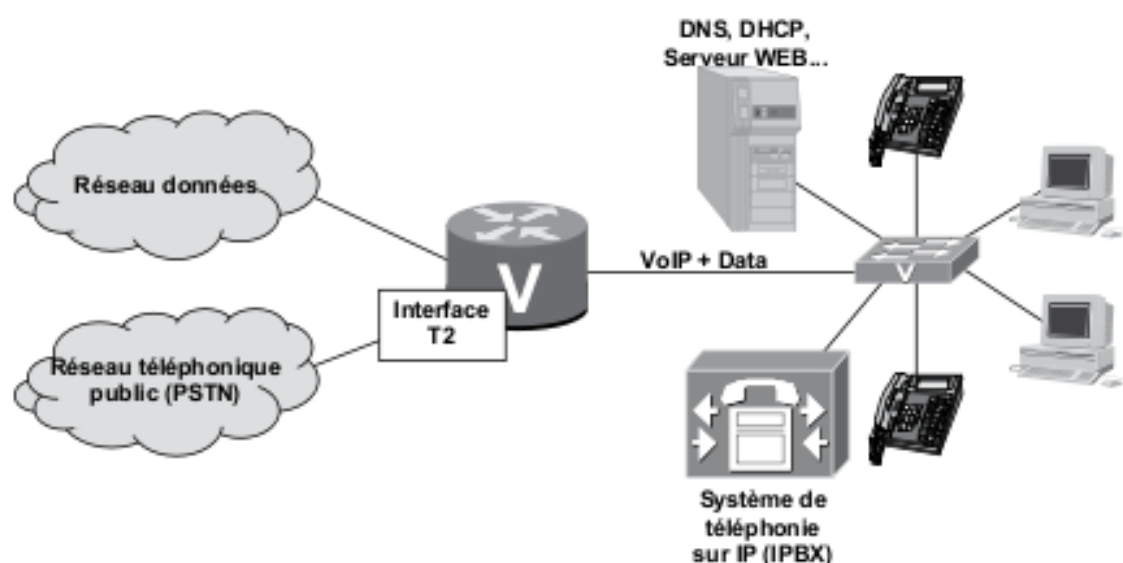


Figure 33.1 Principe de base de la ToIP.

est géré par ce qui est dénommé IPBX. L'installation locale est raccordée au réseau public (PSTN, *Public Switched Telephon Network*) via un routeur doté d'une carte T2 (passerelle voix). L'IPBX est lui équipé d'une interface LAN, les téléphones (téléphones IP) sont directement connectés au réseau local. Le système est dit *téléphonie sur IP* ou *ToIP*.

Deux types de terminaux téléphoniques IP peuvent être distingués :

- ▶ Les « *hardphones* » (*IPPhone* ou *EtherPhone*) correspondent dans l'environnement ToIP au téléphone traditionnel, ils peuvent être dotés d'un mini-écran donnant accès à différents services et en particulier à la messagerie. Dotés d'une « intelligence locale », ils autorisent des fonctions avancées qui vont du simple annuaire téléphonique à un mini-navigateur qui donne accès à l'intranet de l'entreprise voire à Internet. Ces téléphones sont équipés d'un mini-commutateur permettant le raccordement d'un PC. Le poste de travail et le téléphone IP se partageant alors la même connexion physique au réseau ;
- ▶ Les « *softphones* » sont constitués d'un ensemble de logiciels émulant, sur un PC traditionnel équipé d'un microphone et d'un écouteur, un terminal téléphonique. Ce sont les systèmes les plus aboutis en termes de convergence.

Outre les postes IP (*hardphones* ou *softphones*) les systèmes de ToIP doivent pouvoir accueillir les postes téléphoniques analogiques qui sont souvent présents dans les systèmes de sécurité (postes ascenseurs, portiers ou tout autre système à contacts secs reposant sur une connexion téléphonique). Ces différents éléments seront reliés au système de téléphonie via des boîtiers d'adaptation numérique/analogique (ATA, Adaptateur pour terminal analogique).

33.2 La téléphonie, une application parmi d'autres ?

La tendance actuelle est de considérer la téléphonie comme une application parmi les autres mais soumise, vis-à-vis de l'utilisateur, à une contrainte forte : lui garantir une qualité de service comparable à la

téléphonie traditionnelle en terme de qualité, de simplicité¹ et de disponibilité. En téléphonie traditionnelle (figure 33.2), chaque téléphone est relié individuellement au PABX par un réseau dédié, c'est une application spécifique. En téléphonie sur IP, le gestionnaire des appels réside sur un serveur informatique de type micro-ordinateur (PC), les terminaux sont raccordés directement au LAN, c'est une application informatique parmi les autres. Seules des cartes d'interfaçage avec le réseau public, encore souvent en desserte TDM (RNIS), restent spécifiques.

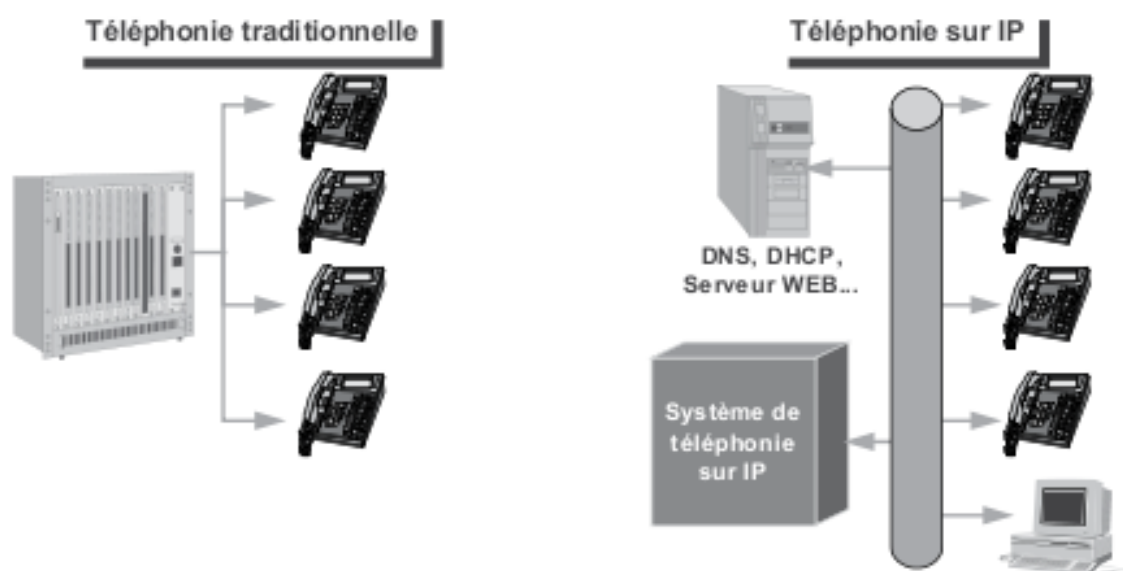


Figure 33.2 Les architectures locales.

33.2.1 Du PABX à l'IPBX

■ Introduction

La voix sur IP dissocie les fonctions traditionnelles d'un PABX en systèmes indépendants (figure 33.3). La voix est traitée par un système dédié au traitement de celle-ci : la *Media Gateway* (passerelle voix). La signalisation est interprétée par un système dévolu à son traitement et à son

¹ « J'ai toujours rêvé d'un ordinateur qui soit aussi facile à utiliser qu'un téléphone... Mon rêve s'est réalisé, je ne sais plus utiliser mon téléphone. » (Bjarn Soustrup, créateur du C++).

éventuelle conversion (*Signaling Gateway*). Enfin, la gestion des communications est assurée par un contrôleur de communication (Gestionnaire d'appels, *Call Manager*, *Call Server*, *Media gateway controller*... les appellations diffèrent selon le constructeur mais les fonctionnalités sont similaires¹) alors que la fonction de commutation est reportée sur le réseau (LAN/WAN). C'est cette dernière fonction qui impose des contraintes sévères au niveau du réseau aussi bien sur le LAN que sur le WAN.

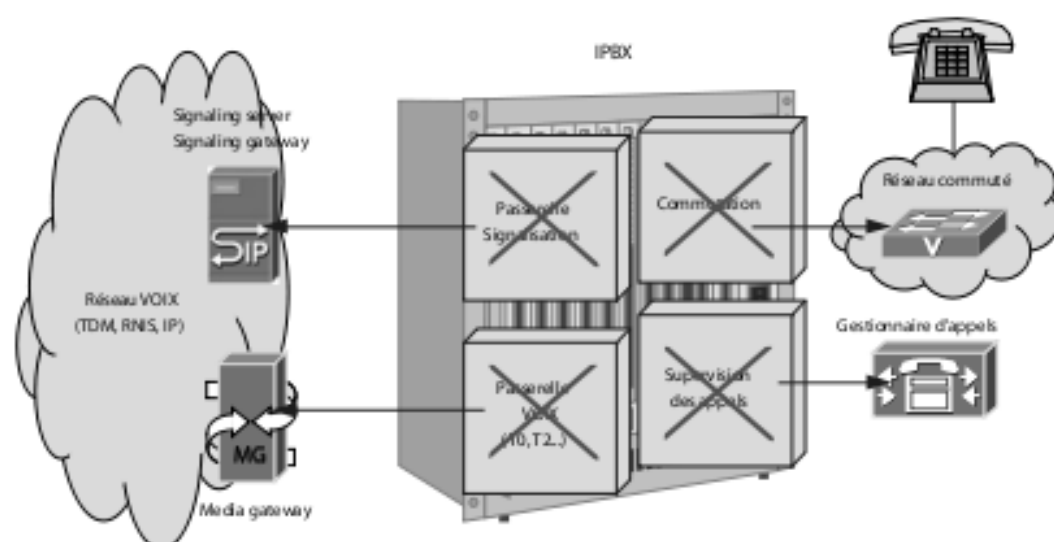


Figure 33.3 Du PABX à l'IPBX.

Un IPBX est un élément qui rassemble en un seul ensemble cohérent toutes les fonctionnalités d'un système de téléphonie IP.

33.2.2 Mécanismes élémentaires d'établissement d'appel

Le schéma simplifié de la figure 33.4 illustre succinctement les différentes phases d'un appel en ToIP (appel local). Les différentes phases d'établissement d'appel sont :

- ▶ l'appelant décroche et numérote, les données d'appel sont transmises au gestionnaire d'appels (*Call Setup*) ;

1 Ces différents termes seront employés indifféremment dans ce chapitre.

- ▶ le gestionnaire d'appels met en relation le numéro appelé et l'adresse IP de l'appelé ; au besoin, il vérifie les droits de l'appelant et de l'appelé (*E.164 Lookup*) ;
- ▶ les données d'appels (*Call Setup*) sont transmises à l'appelé (numéro d'appelant, codec...) ;
- ▶ transmission de l'indication d'appel à l'appelé (sonnerie, *Ring*) ;
- ▶ retour de sonnerie pour l'appelant (*Ring Back*) ;
- ▶ l'appelé accepte l'appel, ce qui provoque directement l'établissement du canal media (7, flux RTP) entre les correspondants.

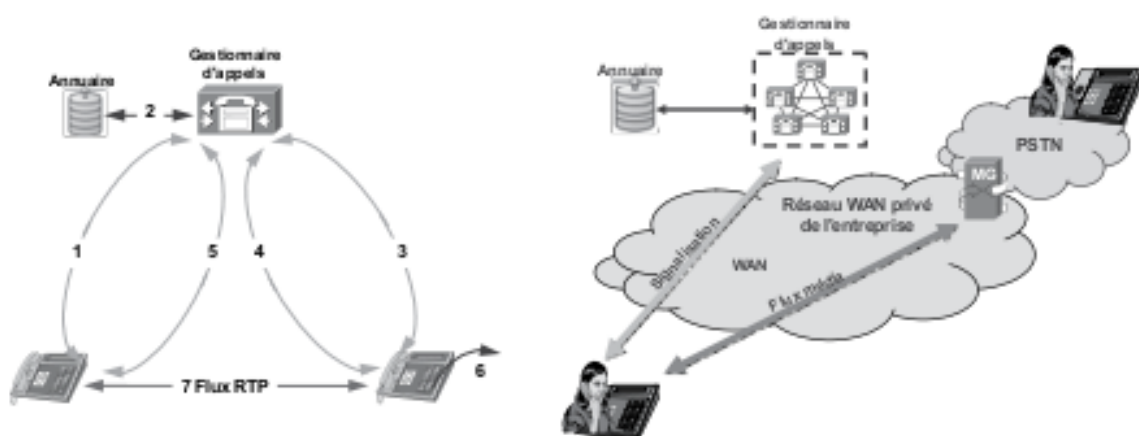


Figure 33.4 Principe d'établissement d'une communication.

Seuls les flux de signalisation (étapes 1 à 5), transitent par le gestionnaire d'appels. Le flux média (7) est établi en mode point à point directement entre les correspondants. Cette indépendance des flux qui dissocie complètement le flux de signalisation du flux media autorise la délocalisation du gestionnaire d'appels (Contrôleur de communications, *Call Server...*), voire de la passerelle *Media Gateway*, vis-à-vis de l'installation locale des terminaux voix.

33.2.3 L'aspect protocolaire

Les mécanismes décrits précédemment mettent en évidence 2 flux : un flux de signalisation et un flux media. Ces deux flux seront transportés par IP (VoIP), protocole en mode datagramme, qui n'assure pas : la

garantie de délivrance, le contrôle d'erreur, le contrôle de séquençement, le contrôle de flux et de congestion ; cependant, il offre des mécanismes de qualité de service permettant d'offrir à chaque type de flux les services réseau dont il a besoin (DiffServ).

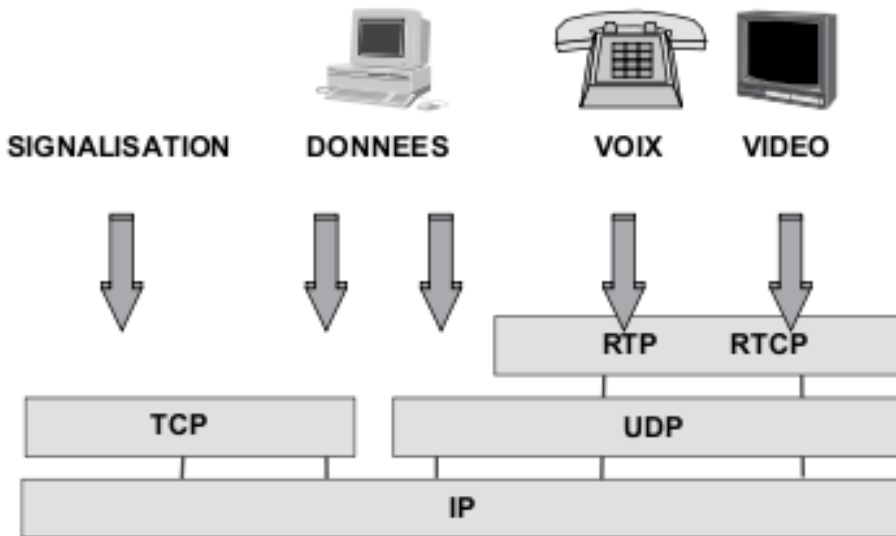


Figure 33.5 Principe de base de la pile protocolaire ToIP.

Au niveau transport, deux possibilités s'offrent : le transport sous TCP ou UDP. Le transport sur TCP est fiable, il établit une connexion de transport, réalise le contrôle et la reprise sur erreur (inadapté pour le flux media), il garantit le séquençement mais réalise un contrôle de flux et de congestion (inadapté pour le flux media). Le transport UDP en mode datagramme ne réalise pas de reprise sur erreur, ni contrôle de flux, ni contrôle de congestion. Il conviendrait pour les flux media mais il ne garantit pas le séquençement. Aussi, pour le flux de signalisation, on optera pour un transport fiable sur TCP¹, alors que le flux media nécessitera un protocole additionnel au-dessus d'UDP pour garantir le séquençement. Le protocole RTP (*Real Time Protocol*) horodate les datagrammes et permet un contrôle de séquençement, on lui adjoindra un protocole complémentaire fournissant un rapport sur la qualité du transport : RTCP (*Real Time Control Protocol*). La pile de base est représentée figure 33.5.

1 Bien que SIP soit implémenté aussi sur UDP.

34

L'architecture logique et la signalisation

Une architecture voix sur IP, outre les téléphones IP, est organisée autour d'une passerelle voix/vidéo (*Voice/Video Gateway*) et d'un contrôleur de communication (gestionnaire d'appels, *Call manager*...). La passerelle (figure 34.1) réalise l'interconnexion de l'installation voix locale vers les réseaux traditionnels : réseau en mode paquet (WAN IP) et les réseaux en mode circuit (RTC et RNIS).

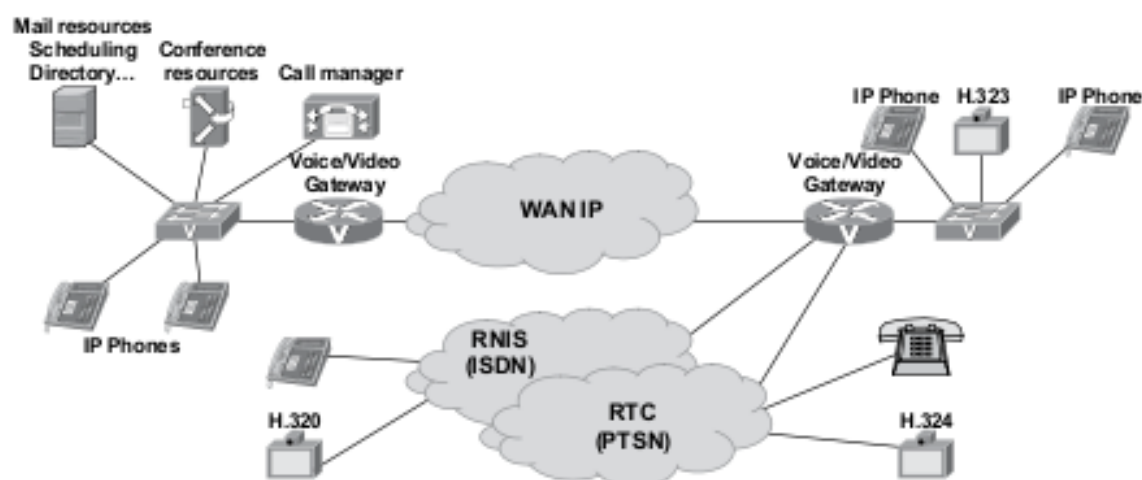


Figure 34.1 L'architecture ToIP de principe.

Elle assure aussi les conversions nécessaires de format des flux voix/vidéo et la translation des protocoles de signalisation. Concrètement, une passerelle voix peut être un PABX traditionnel auquel on a ajouté une interface IP ou un routeur équipé de cartes voix. Le contrôleur de communication réalise le contrôle des appels, la mise en correspondance d'un numéro de téléphone avec une adresse IP et le routage des appels distants vers la passerelle. Un troisième élément, optionnel, gère l'établissement, le mixage et la diffusion des conférences (*Conference resources*). Ces différentes fonc-

tionnalités peuvent être localisées dans une même entité (IPBX) ou réparties en plusieurs éléments.

34.1 L'architecture H.323 de l'UIT-T

La recommandation H.323 définit un modèle architectural pour assurer le transport de la voix sur un réseau avec ou sans qualité de service. L'architecture H.323 comprend diverses fonctionnalités (ou éléments) représentées figure 34.2.

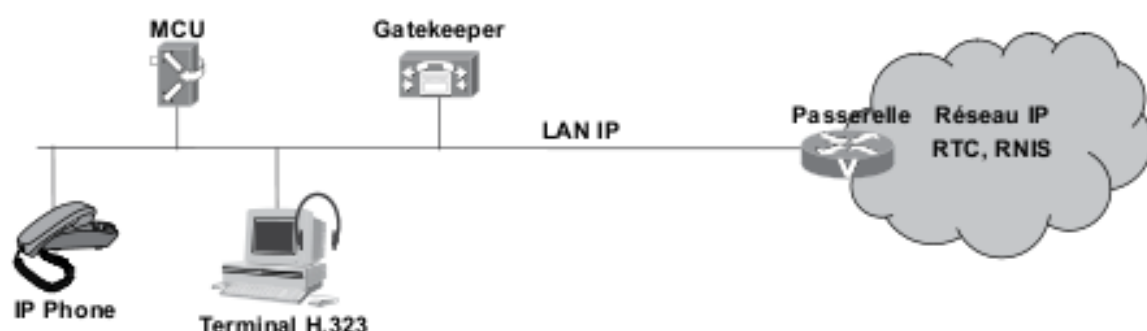


Figure 34.2 L'architecture matérielle d'une zone H.323.

Les terminaux H.323 sont raccordés directement au LAN IP. Ils ont la capacité d'établir des communications voix, vidéo et/ou données en temps réel avec tout terminal de la zone H.323 ou non, en mode point à point, multipoint ou diffusion. L'appel est réalisé selon le protocole Q.931 (protocole D du RNIS). Les protocoles mis en œuvre par un terminal H.323 sont (figure 34.3) :

- ▶ H.225 ou **RAS** (*Registration Admission Status*), ce protocole gère l'enregistrement du poste terminal auprès d'une passerelle (*Registration*), il émet une demande de ressource auprès du *gatekeeper* (*Admission* et *Status*). Il est également chargé de la signalisation et de l'établissement d'un appel (sous-ensemble du protocole Q.931 du RNIS) ;
- ▶ H.245, ce protocole permet aux terminaux d'échanger leurs capacités audio/vidéo (codecs supportés, nombre de canaux possibles, modes de conférence acceptés...) et de négocier les canaux logiques de dialogue ;

- T.120, ce protocole optionnel gère l'échange de données entre terminaux H.323.

La figure 34.3 décrit la pile protocolaire H.323. La voix est transportée en mode datagramme sur UDP tandis que la signalisation est transportée en mode connecté sur TCP (encapsulation **TPKP**, *Transport PacKeT*, RFC 1006). Les spécifications H.323 correspondent aux niveaux session et supérieurs du modèle de référence, cette approche assure l'interopérabilité des systèmes quel que soit le réseau de transport utilisé.

La passerelle H.323 ou *Voice/Video Gateway* assure l'interface avec une entité H.323 et une entité non H.323 comme les réseaux RNIS (H.320) ou ATM (H.321), la conversion de signalisation H.225/Q.931, l'adaptation des supports et des débits. Chaque passerelle H.323 connaît les numéros E.164 (numéros de téléphone) qui lui sont rattachés, elle dispose en mémoire d'une table de correspondance qui associe à un numéro E.164 une adresse IP, un e-mail ou un alias. Si le réseau est important, la maintenance des tables peut devenir vite impossible. Ce problème trouve

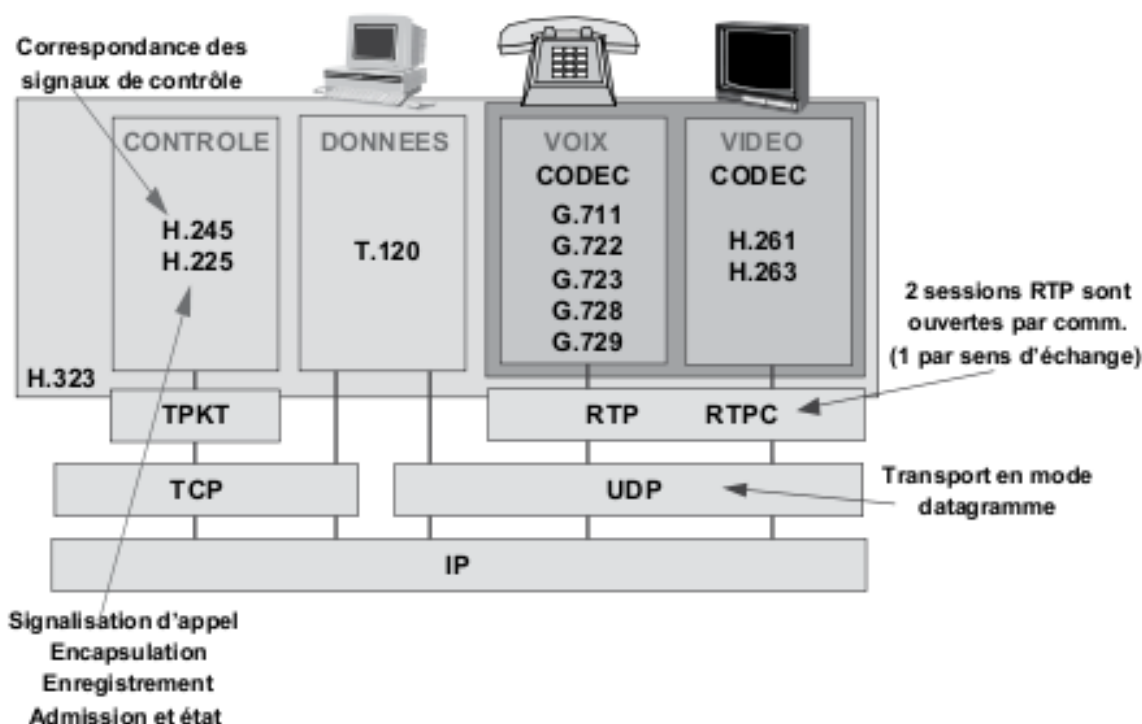


Figure 34.3 La pile protocolaire H323.

sa solution par l'emploi d'un *gatekeeper*¹ qui va centraliser les tables de conversion d'adresses. Chaque *gateway* vient s'enregistrer sur son *gatekeeper* et lui déclare toutes ses adresses E.164. Lorsqu'une passerelle doit établir un appel, elle s'adresse au *gatekeeper* qui lui fournit l'adresse IP de la passerelle destination.

Le portier H.323 ou *gatekeeper*, système optionnel de gestion des communications établies par les entités H.323 fournit les services :

- ▶ d'enregistrement (*Registration*) en établissant une liste des utilisateurs joignables sur un terminal donné et effectue la translation d'adresses (alias, e-mail, E.164...);
- ▶ d'admission des appels en contrôlant les droits des utilisateurs (rejet éventuel d'appel) et en gérant la bande passante;
- ▶ de gestion de la passerelle H.323 (management de la zone);
- ▶ de journalisation des appels, taxation.

Enfin, le **MCU** (*Multipoint Control Unit*), équipement optionnel, gère l'établissement, le mixage et la diffusion des conférences (contrôle des liaisons multipoints en *multicast*).

La figure 34.4 illustre les relations protocolaires entre un terminal H.323 et un équipement non H.323, dans ce cas de figure la liaison entre le PABX et la passerelle voix est généralement réalisée par des cartes voix présen-

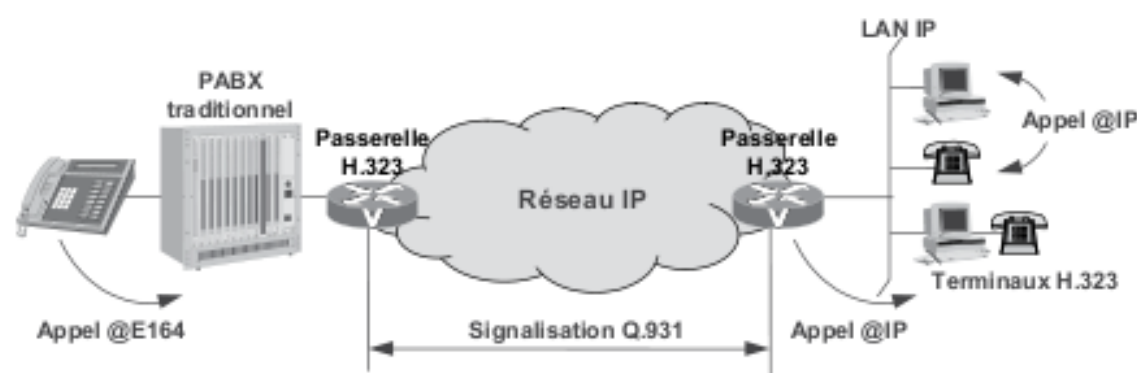


Figure 34.4 La mise en relation selon H.323.

1 Littéralement : garde-barrière ; certains utilisent le terme de portier, ce qui est réaliste en terme de fonctionnalités. La plupart des ouvrages en langue française utilise le terme anglo-saxon, nous nous conformerons à cet usage.

tant une interface de type MIC, E1 ... Le terminal téléphonique, non H.323, établit un appel en direction du PC multimédia, le routeur (*Gateway* H.323) interprète la numérotation et initialise un appel Q.931 vers l'agent H.323 distant. L'agent H.323 réalise la correspondance entre une adresse E.164 (Q.931) et une adresse IP, il établit un canal de communication entre le demandé et le demandeur.

34.2 Le protocole SIP de l'IETF (RFC 3261)

34.2.1 Généralités

Actuellement, la plupart des solutions développées n'utilisent plus la signalisation H.323 (v1, v2 ou v3) d'origine UIT-T que pour des problèmes de compatibilité avec l'existant. H.323 a définitivement cédé la place à SIP. Développé au sein du groupe de travail **MMUSIC** (*Multiparty Multimedia Session Control*) de l'IETF, le protocole **SIP** (*Session Initiation Protocol*) est beaucoup plus simple que H.323.

Inspirés par le protocole HTTP (*Hyper Text Transfer Protocol*), les messages SIP sont au format texte, ce qui confère au protocole une grande évolutivité. L'architecture est du type client/serveur, un échange forme une transaction, il est composé d'une requête émise par le client et d'une réponse fournie par un serveur. La figure 34.5 illustre une architecture SIP simplifiée.

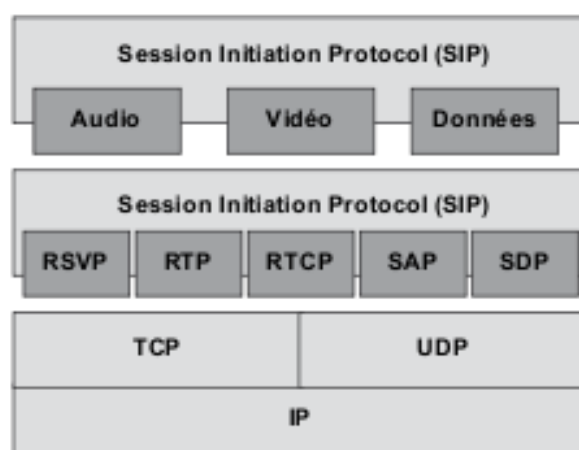


Figure 34.5 L'architecture protocolaire SIP.

À l'instar de H.323, SIP s'appuie sur les protocoles temps réel (RTP et RTCP), il peut éventuellement utiliser RSVP pour obtenir une certaine qualité de service sur le réseau. SIP ne fait qu'initialiser une session, il lui est adjoint de nombreux protocoles complémentaires :

- ▶ le protocole **SDP** (*Session Description Protocol*, RFC 2327) fournit la description des sessions multimédia :
 - formatage des messages ;
 - gestion des sessions (nom, date, objet...) ;
 - description des flux (audio, vidéo...) ;
 - paramètres des flux (adresses, ports, formats...) ;
- ▶ le protocole **SAP** (*Session Announcement Protocol*) informe de l'ouverture d'une session multimédia en mode *multicast* ou non ;
- ▶ le protocole **SCCP** (*Single Conference Control Protocol*) ;
- ▶ le protocole **RTSP** (*Real Time Streaming Protocol*, RFC 2326) qui définit :
 - la description des media échangés (codage, nature...) ;
 - le choix des ports ;
 - et introduit la notion de session au sens ISO du terme.

Basé sur le modèle client serveur, SIP distingue deux types d'agent : les clients et les serveurs. Les clients ou **UAC** (*User Agent Client*) sont les équipements à l'origine des appels SIP (téléphones IP) ou des passerelles voix. Les passerelles voix SIP ont les mêmes fonctionnalités que les passerelles H.323.

Les agents serveurs (**UAS**, *User Agent Server*) sont des équipements classiques (serveur Windows, Linux...) qui regroupent les services offerts par SIP. Ce sont :

- ▶ les serveurs d'enregistrement utilisés pour la localisation des utilisateurs (*Registrar* ou *Location Server*). Les serveurs d'enregistrement contiennent toutes les caractéristiques des agents SIP autres que les passerelles, ils gèrent les requêtes *Register* envoyées par les *users agents* pour notamment assurer l'unicité des URI (*Uniforme Ressource Identifier*) qui identifient un utilisateur, la syntaxe est :
 - ▶ « sip :nom@domaine.com » ;

- les serveurs de délégation (*Proxy server*) qui gèrent les clients SIP, reçoivent et transmettent les requêtes au serveur suivant (*Next-hop server*). Le SIP *Proxy* a un rôle similaire au *gatekeeper* de H.323 (*Call manager* ou contrôleur de communication). Un SIP *Proxy* peut interroger un SIP *Registrar* ou un DNS pour acquérir les informations d'acheminement de la signalisation et des communications ;
- les serveurs de redirection (*Redirect server*) qui, sur requête et après une éventuelle consultation du *Registrar*, transmettent l'adresse du *next-hop server* à l'agent client.

La figure 34.6 illustre un appel SIP vers un utilisateur qui s'est déplacé. Le client appelant envoie une requête INVITE au serveur proxy auquel il est relié. Ce message contient l'adresse connue du destinataire. Le *proxy server* interroge le *location server* (DNS, LDAP ou autre) qui lui fournit la nouvelle adresse, le *proxy* redirige la requête vers la nouvelle adresse de l'appelé (INVITE). Le poste appelé sonne et le poste appelant reçoit un message de retour de sonnerie (SIP 180). L'appelé décroche signifiant

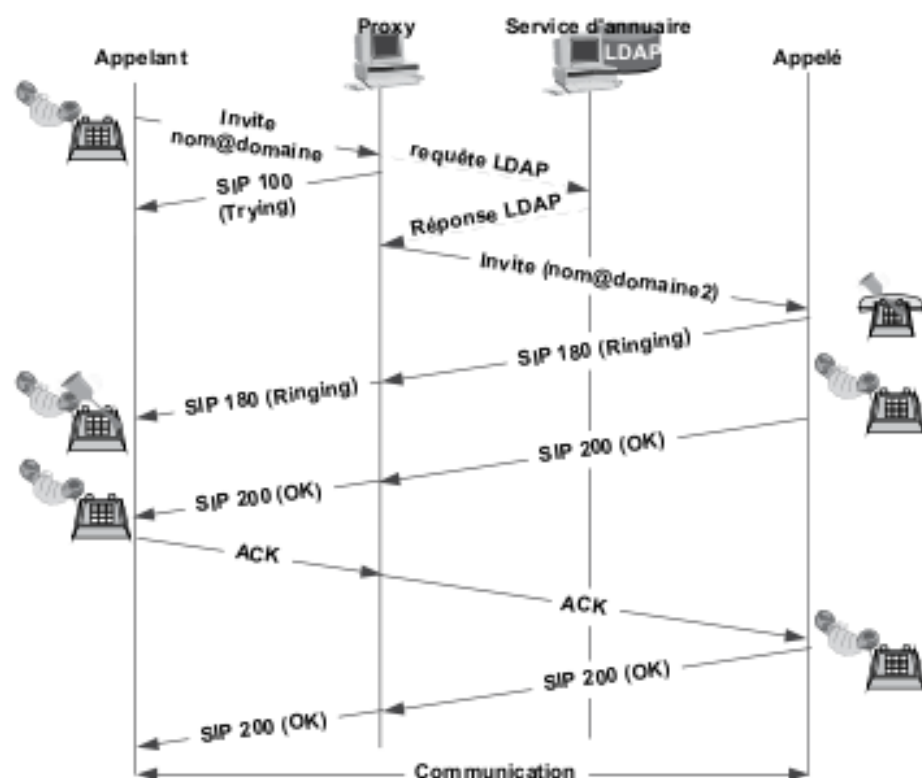


Figure 34.6 L'ouverture d'une communication SIP en mode proxy.

ainsi son acceptation de la communication, le système émet alors un message SIP 200 (OK). L'appelant acquitte le message d'acceptation. Le message BYE acquitté met fin à la communication.

Sans *Server Proxy*, l'appelant se serait adressé à un *redirect server*, celui-ci consulte le *location server* (serveur DNS, LDAP...) et envoie à l'appelant la nouvelle adresse dans un message **SIP 302** (*Moved temporarily*) ou **SIP 301** (*Moved permanently*). L'appelant émet alors une requête INVITE directement à l'appelé.

34.2.2 Exemple de message SIP

La syntaxe SIP utilise la même syntaxe que les messages http/1.1 (RFC 2616) et le jeu de caractères (ISO 10646 et codage UTF-8 du RFC 2279). Deux messages dont la structure est illustrée ci-après sont utilisés par SIP : les requêtes et les réponses. Les tableaux 34.1 et 34.2 fournissent un exemple de message SIP.

Tableau 34.1 Structure générale d'un message SIP.

Requête (client vers serveur)	Réponse (serveur vers client)
Ligne de requête (Méthode, Requête URI, version SIP)	Ligne d'état (Version SIP, code d'état, motif)
En-tête général	En-tête général
Caractères CR/LF délimitent la fin du champ en-tête du corps du message	Caractères CR/LF délimitent la fin du champ en-tête du corps du message
Corps du message	Corps du message

Tableau 34.2 Exemple de message SIP/SDP, appel d'Alice à Bob.



Ligne du message	Signification
INVITE sip : bob@domaine2.com SIP/2.0	INVITE est le nom de la méthode invoquée.
Via : SIP/2.0/UDP 10.0.0.1 ; branch=azertyui	Contient l'adresse à laquelle Alice attend les réponses à sa requête.

Ligne du message	Signification
MAX-Forwards :70	Limite le nombre de sauts entre la source et la destination, Max-Forwards est décrémenté de 1 à chaque saut.
To : Bob<sip : bob@domaine2.com>	Contient le nom d'affichage (Bob) ainsi que l'URI SIP de Bob.
From : Alice<sip.alice@domaine1.com>	Contient le nom d'affichage d'Alice et l'URI SIP de l'origine de la demande.
Call-ID : wx9Bcvbn@10.0.0.1	Est un identifiant d'appel, unique au monde pour cet appel. L'association de l'étiquette To, From et du Call-ID définit une relation SIP d'homologue à homologue entre Alice et Bob dénommée « dialogue ».
CSeq :123456 INVITE	Ou Command Sequence, contient un entier et un nom de méthode, le nombre Cseq est incrémenté à chaque nouvelle demande.
Contact : <sip : alice@10.0.0.1>	Contient un URI SIP (sip :alice@domaine1.com) pour joindre Alice, voire l'adresse IP d'Alice.
Content-Type : application/sdp	Contient la description du corps de message.
Content-Length :XXX	Indique la longueur en octets du corps de message.
CR/LF	Ligne séparant l'en-tête du corps de message.
V = 0	SDP version du protocole, 0=version mineure la seule définie aujourd'hui soit la 1.
o = Alice 12345678 12345679 IN IP4 10.0.0.1	Origine de l'appel, identificateur de session, version de session, réseau (IN IP4) et origine de l'appel.
c = IN IP4	Type de réseau.
T = 0 0	Durée de la session (0 0) session non limitée dans le temps (début 0 : début non défini, fin 0 fin indéterminée).
m = audio 45120 RTP/AVP 0	Type de flux media invoqué, port origine, protocole RTP sur UDP (AVP), l'encodage est précisé ligne suivante.
a = rtpmap :0 PCM/8 000	Encodage PCM à 8 kbit/s.

Le téléphone d'Alice ne connaissant pas la localisation de Bob (172.16.0.1) émet un appel en direction du serveur *proxy* qui dessert son domaine, le serveur Proxy transmet la requête au nom d'Alice, il ajoute une ligne « Via » au message pour indiquer que les messages retour devront transiter par lui. De même, il transmet la requête au serveur *proxy* du domaine de Bob qui ajoutera une ligne « Via » au message.

34.3 Signalisation, la synthèse

Pendant longtemps, H.323 a dominé, issu de l'IUT, ce protocole robuste et complet est cependant complexe. SIP se présentait comme outsider, car insuffisamment développé, il n'offrait que peu de services. Ces protocoles ont convergé : meilleure efficacité de H.323, renforcement des services de SIP. Aujourd'hui SIP domine dans les systèmes ToIP. Cependant MGCP/MEGACO//H.248 reste le protocole dominant dans les systèmes de VoIP pour la commande des passerelles d'interconnexion. MGCP est aussi le seul protocole permettant le pilotage des postes analogiques.

35.1 L'architecture générale

35.1.1 Généralités

La ToIP consiste à déporter la fonction de commutation sur l'intégralité du réseau. Le LAN et le WAN peuvent alors être considérés comme la matrice de commutation du système IPBX. L'une des conséquences de cette approche est qu'il est possible de localiser en un point l'intelligence (gestionnaire d'appels, *Call Manager*...) en un autre les passerelles (*Media Gateway*...) et les terminaux. Cette faculté a donné naissance à une architecture dite Centrex. Un Centrex peut faire l'objet d'une offre opérateur (externalisation de la téléphonie) ou constituer la base de l'architecture d'un réseau privé (Centrex privé), c'est généralement le type d'architecture adoptée par les sociétés multisites.

35.1.2 L'architecture centralisée (Centrex privé)

Le Centrex peut mutualiser les accès au PSTN (Réseau public de téléphonie) en un seul point. Si tous les sites sont implantés dans la même zone de taxation, cette approche peut être porteuse d'économie, elle facilite l'administration et permet de voir le réseau comme une entité unique (même séquence SDA¹) mais elle fragilise le réseau en cas de défaillance

1 Un numéro SDA (Sélection directe à l'arrivée) permet de mettre en relation directe un poste téléphonique extérieur appelant avec un poste téléphonique de l'entreprise sans passer par un standard.

de la *Media Gateway* et du *Call Manager* (gestionnaire d'appels, *Call Server*...).

Pour pallier la fragilisation du système dû à la concentration des accès au réseau public, chaque site important de la figure 35.1 est doté d'un accès au réseau public. Dans cette approche les appels, vers le réseau public, peuvent être établis localement ou décentralisés. En cas de défaillance d'un attachement local les appels peuvent être réacheminés vers un autre attachement (survivabilité). De même, en cas de saturation du réseau, les appels *On-Net* peuvent être établis en débordement sur le réseau public.

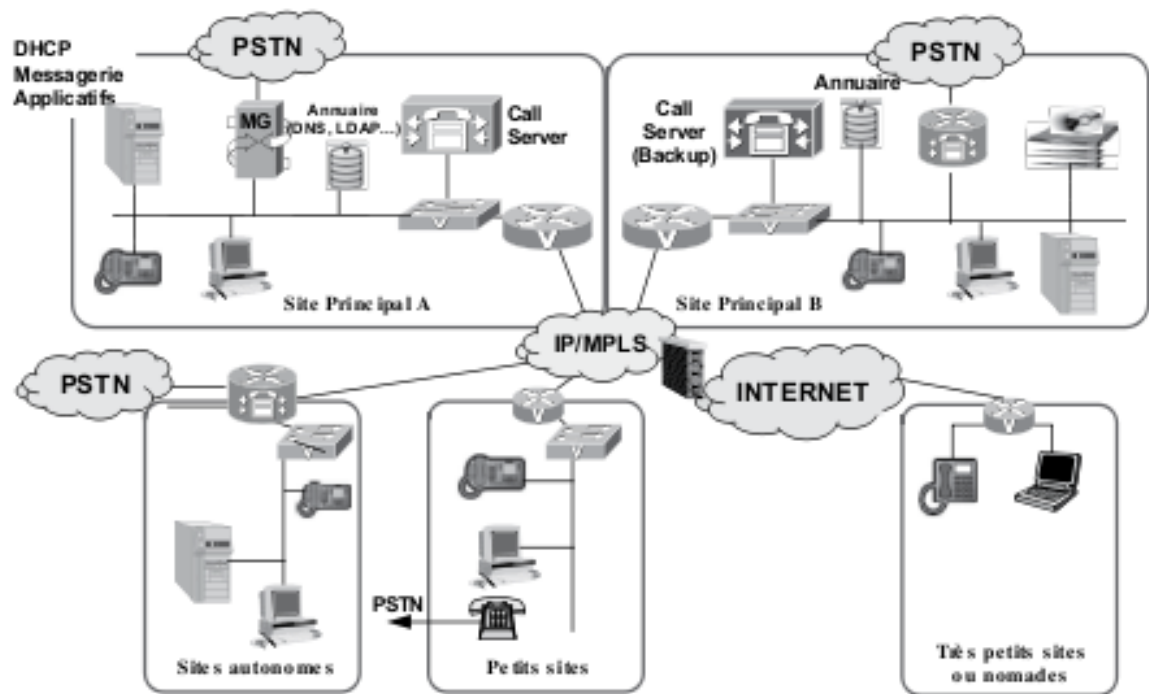


Figure 35.1 Centralisation de la gestion avec passerelle locale.

Dans cette architecture, les petits sites peuvent ne pas être dotés d'accès au réseau public. Cependant, pour des raisons de sécurité, il est indispensable que ces derniers soient dotés d'un abonnement de secours, non relié au réseau et qui éventuellement peut servir de ligne Fax au site considéré. Les très petits sites (travailleurs distants ou à domicile) peuvent accéder aux services téléphoniques de l'entreprise, ils sont raccordés comme les nomades *via* Internet ou autre avec éventuellement un protocole d'authentification (Radius...).

35.1.3 Notion de survivabilité

Dans cette architecture, courante dans les entreprises, le gestionnaire d'appels (*Call Manager, Call Server...*) détient l'intelligence du réseau, il est indispensable qu'il soit secouru (duplication). Cependant, l'implantation d'un site de *backup* pour une reprise des communications en temps réel est soumise à des contraintes drastiques en termes de délai d'acheminement et de débit entre le site principal et le site de *backup*. Ces contraintes sont liées à l'architecture du constructeur, elles varient énormément d'un constructeur à un autre (de quelques dizaines de millisecondes à plusieurs secondes).

Si, malgré la centralisation de la gestion, chaque site dispose d'un raccordement au réseau public, les appels *on-net* ne pouvant plus être acheminés par le réseau, le seront par le réseau public durant le temps de rétablissement du réseau. Cette option garantit la permanence du service, elle implique, sur chaque site, l'autorisation de débordement et nécessite que la passerelle locale abrite un mini-gestionnaire d'appel (logiciel de survivabilité).

La survivabilité (figure 35.2) concerne tous les mécanismes qui permettent à un site, en cas de défaillance du gestionnaire d'appels ou du réseau, de continuer à avoir accès aux services téléphoniques. Lors du fonctionnement normal, un poste téléphonique s'enregistre auprès du serveur d'appels (1). Périodiquement, le poste émet des messages « *keepalive* » (2), le serveur d'appels répond qu'il est actif (3). En cas de non-réponse, le poste s'enregistre alors auprès de la passerelle locale (*media gateway*). La signalisation (établissement d'appels...) n'est plus acheminée vers le gestionnaire central, elle est traitée localement. Dès la reprise de l'activité du serveur d'appel (réponse aux messages « *keepalive* »), le poste se désinscrit.

Ce mécanisme n'impacte pas les communications en cours qui restent maintenues (le flux média est indépendant du flux de signalisation). Généralement durant la période de survivabilité, les services téléphoniques rendus le sont à minima.

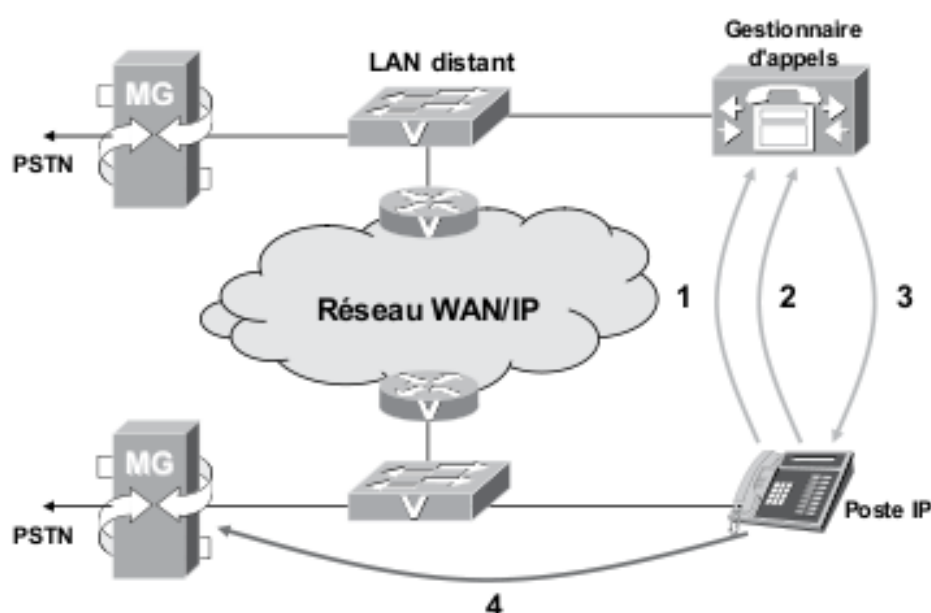


Figure 35.2 Principe de base du mécanisme de survivabilité.

35.1.4 L'annuaire d'entreprise et l'annuaire téléphonique

L'annuaire d'entreprise est devenu une composante essentielle du système d'information des grandes entreprises. La notion d'annuaire unifié pour séduisante qu'elle soit, n'est pas une réalité concrète.

Quels que soient le système et les constructeurs, les services téléphoniques font appel à un annuaire spécifique généralement embarqué dans le système de téléphonie. Cette approche se justifie d'une part par le nombre de sollicitations du système et d'autre part par la nécessité de la permanence du service téléphonique même en cas de défaillance de l'annuaire d'entreprise.

En téléphonie traditionnelle, les annuaires (entreprise et téléphonique) sont indépendants les uns des autres et sources de nombreuses incohérences. En téléphonie sur IP, l'annuaire d'entreprise peut être enrichi pour contenir certaines informations complémentaires relatives à la téléphonie et l'annuaire téléphonique initialisé à partir de l'annuaire d'entreprise. Les utilisateurs sont créés ou détruits uniquement en un seul point : l'annuaire d'entreprise. Les informations relatives à l'utilisateur sont gérées par l'annuaire d'entreprise, l'annuaire téléphonique n'est enrichi

que des données spécifiques à la téléphonie (facilités téléphoniques, droits spécifiques à un utilisateur.).

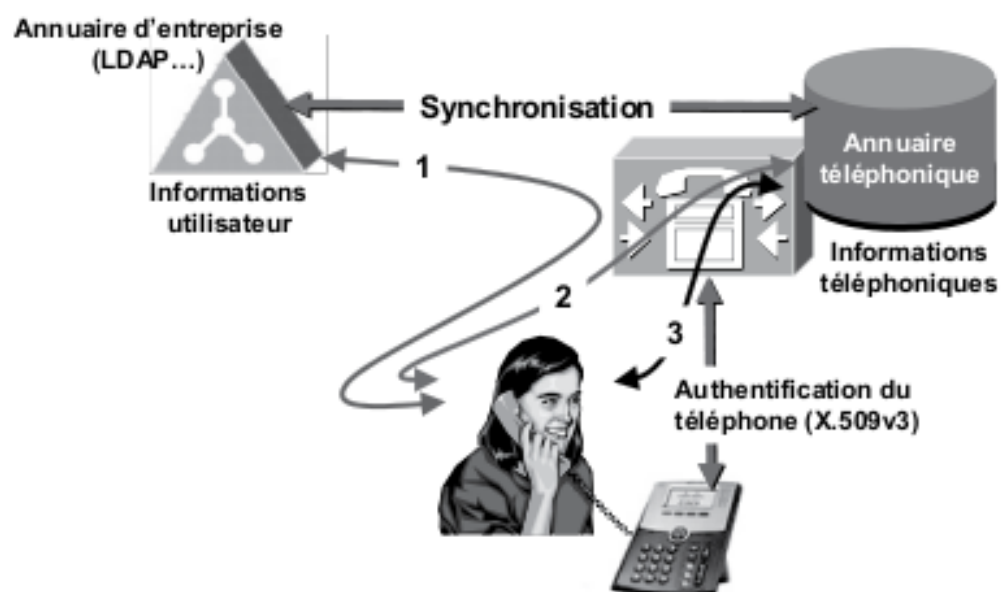


Figure 35.3 Relations entre les annuaires.

La figure 35.3 illustre les relations entre les deux annuaires. Le gestionnaire d'appel peut contenir, en cache, toutes les données utilisateurs. Lors de la connexion d'un utilisateur, si ses données ne sont pas contenues dans le cache du gestionnaire d'appels, l'utilisateur est authentifié par l'annuaire d'entreprise et ses données sont mises en cache pour une utilisation ultérieure (repères 1 et 2), il en est de même des données utilisateur propres à la téléphonie (3, facilités téléphoniques).

35.1.5 Numéros d'urgence

■ Problématique de la localisation géographique

Lors d'un appel d'urgence (E.911 aux États-Unis pour *Enhanced* 911, E.112 en Europe, E.000 en Australie), en téléphonie traditionnelle, la localisation de l'appelant est déduite de son attachement physique au réseau (commutateur public). Lorsque le système d'IPBX est isolé et qu'il dispose de son propre attachement au réseau public de téléphonie, la question ne se pose pas, elle se résout de la même manière qu'en téléphonie traditionnelle.

Ce n'est pas le cas dans les architectures de type « Centrex privé » ne disposant que d'une seule passerelle vers le réseau public (figure 35.4). Dans ce cas, ce seront les secours rattachés géographiquement à la passerelle qui seront alertés !

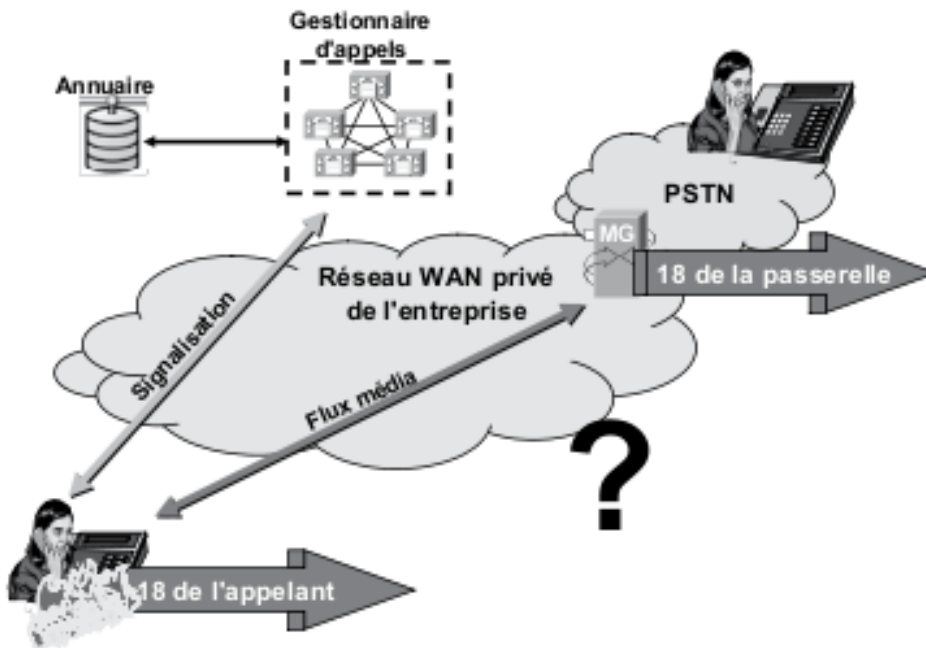


Figure 35.4 Établissement des numéros d'urgence.

Dans ce mode de réalisation, il est indispensable que l'adressage IP soit géographique (chaque site doit pouvoir être localisé), et que l'appel du numéro court soit traduit en « numéro noir¹ » des services locaux à appeler. La mise en œuvre de l'appel d'urgence centralisé se heurte à la mise en relation d'une adresse IP appelant et du « numéro noir » à composer. Compte tenu que l'accès au « Plan départemental des appels d'urgence » est réservé aux opérateurs de téléphonie, l'entreprise n'ayant pas accès à cette liste, cette solution est inenvisageable pour un centrex privé. Cette dernière contrainte impose que chaque site dispose d'une passerelle locale pour assurer, au minimum, l'acheminement des numéros d'urgence.

¹ Le numéro noir correspond au numéro téléphonique réel du service demandé. Localement, ce numéro est mis en correspondance avec le numéro public dit numéro court (15, 17, 19, 112...).

Enfin, pour établir un appel, il est nécessaire que des ressources soient disponibles, que cet appel soit *on-net* ou *off-net*. Aussi, se pose la question de la préemption de ressource et de la priorisation de celle-ci. Ce problème est pris en compte par le protocole **MLPP** (*Multi Level Precedence and Preemption*) d'origine militaire (États-Unis) et qui semble aujourd'hui être rendu obligatoire par les différentes administrations nationales (MLPP ou protocole similaire). La priorité maximale étant affectée aux numéros d'urgence, d'autres niveaux peuvent être définis en fonction de la hiérarchie dans l'entreprise.

■ Accès aux numéros d'urgence et configuration du poste téléphonique

En téléphonie traditionnelle, c'est le poste téléphonique qui dispose de droits (appels extérieurs...) ; en téléphonie sur IP, c'est l'utilisateur qui se voit octroyer des droits. En l'absence d'utilisateur identifié par le système (login), le poste ne dispose d'aucun droit et par conséquent ne peut être utilisé. Aussi, pour permettre à toute personne, même n'appartenant pas à l'entreprise d'établir un numéro d'urgence, un poste téléphonique IP disposera d'un minimum de droits hors utilisateur connecté (profil par défaut), droits qui seront étendus lors de l'identification d'un utilisateur (figure 35.5). Cette indépendance poste physique/droits autorise la mobilité de l'utilisateur qui ainsi peut sur n'importe quel poste de l'entreprise retrouver l'intégralité de ses droits.

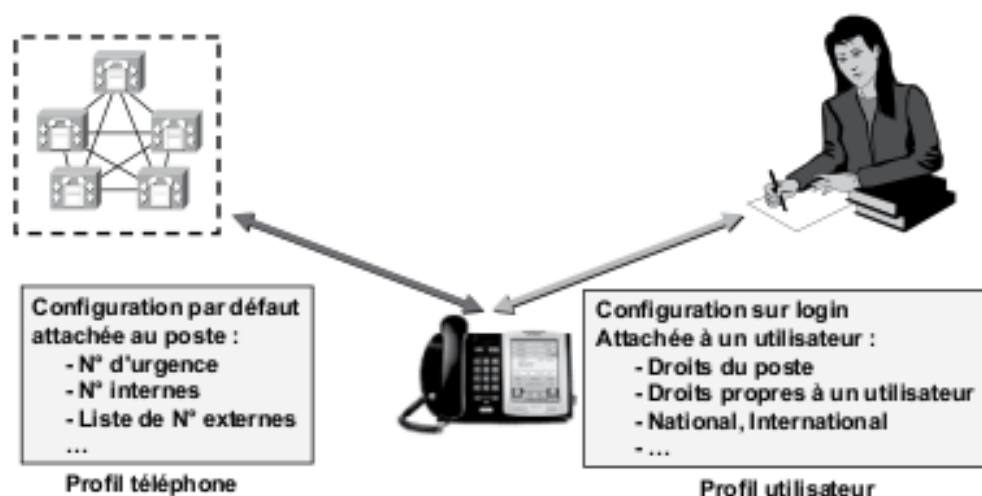


Figure 35.5 Configuration des postes.

35.2 La qualité de service

35.2.1 Généralités

L'approche QoS réseau n'est que l'une des composantes de la qualité de service attendue par l'utilisateur. Au niveau d'un service téléphonique, celle-ci englobe de nombreux autres indicateurs tels que :

- ▶ la qualité de la restitution vocale, souvent le premier facteur d'appréciation de l'outil résulte de la qualité du codec utilisé, de celle du terminal et de celle du réseau (latence, gigue, taux de perte de paquets dans le réseau...) ;
- ▶ la qualité d'obtention et de maintien du service s'exprime par le taux d'accès (ou de refus) au service, temps d'établissement de l'appel, taux d'interruption intempestive de la communication, fiabilité globale du système, taux d'appels perdus.

La qualité d'une communication téléphonique subit deux sources de détérioration : l'une en relation directe avec la qualité de restitution et d'appréciation de la qualité vocale, l'autre sur la notion d'interactivité de l'échange (figure 35.6).

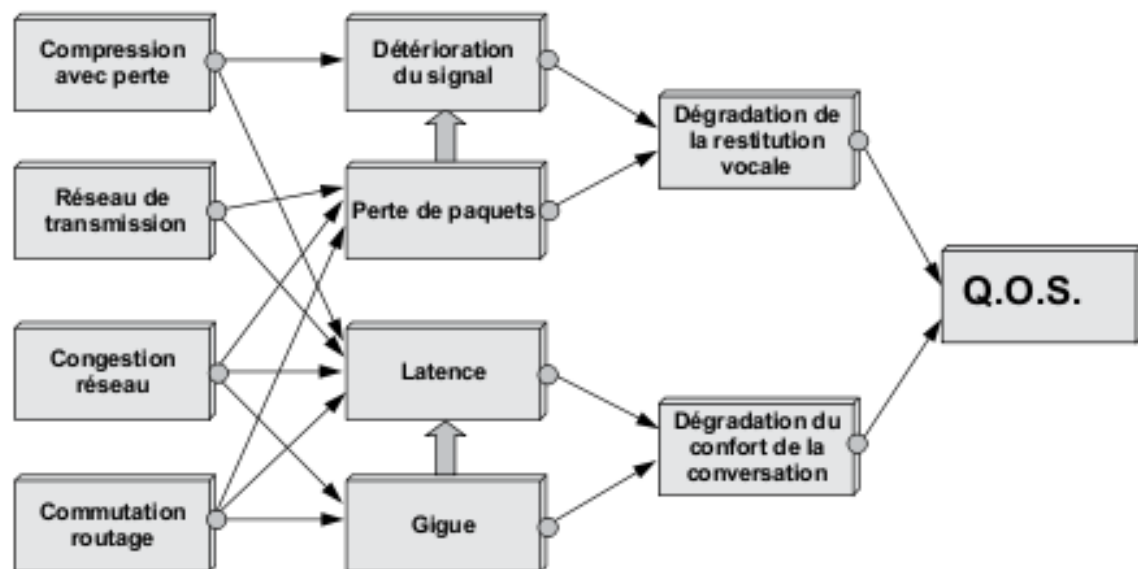


Figure 35.6 Les éléments déterminants de la QoS.

L'oreille est un organe très sensible, capable de distinguer de petites variations de qualité. Les locuteurs perçoivent très rapidement la moindre dégradation (bruit, distorsions, mini-coupures, voix hachée, écho...). Le nouveau service installé doit tendre à offrir un service de qualité équivalente à celle que l'utilisateur connaissait précédemment, alors que le réseau de transport n'est que *best effort* (IP) et que, compte tenu des contraintes de sécurité imposées par la ToIP, l'utilisateur n'est pas obligatoirement réceptif à cette nouvelle technologie.

35.2.2 Appréciation de l'interactivité de la communication

Pour permettre un échange interactif, la voix est transmise sous contraintes de délais. Les chiffres suivants, tirés de la recommandation UIT-T G.114 et donnés à titre indicatif, précisent les classes de qualité et d'interactivité en fonction du délai de transmission dans une conversation téléphonique (tableau 35.1).

Tableau 35.1 Classes de qualité d'une communication selon l'ITU.

Classe	Délai par sens	Commentaires
1	0 à 150 ms	Acceptable pour la plupart des conversations
2	150 à 300 ms	Acceptable pour des communications faiblement interactives Comme pour les communications par satellite (250 ms)
3	300 à 700 ms	Communication pratiquement en <i>half duplex</i>
4	Au-delà de 700 ms	Inutilisable pour des applications de téléphonie (<i>half duplex</i>)

Les sources de délais sont nombreuses dans un système de VoIP, elles sont dues :

- ▶ au codage/décodage et à la mise en paquets de la voix ;
- ▶ à la sérialisation (émission sur le lien local) ;
- ▶ à la gestion des files d'attente des éléments actifs, d'où la nécessité de priorisation des flux (QoS) ;
- ▶ à la latence des commutateurs ;

- ▶ au délai de propagation ou au délai de transfert sur un réseau opérateur (latence réseau) ;
- ▶ à la compensation de gigue.

35.3 Conclusion

Aujourd'hui, la question n'est plus de savoir s'il faut adopter la ToIP. Tous les renouvellements d'équipements téléphoniques optent pour cette solution. La ToIP est arrivée à maturité en termes de services, de qualité et de fiabilité. Bien qu'à ce jour tous les réseaux d'opérateur aient migré en VoIP, la distribution reste diverse (figure 35.7) :

- ▶ la liaison résidentielle utilise toujours un raccordement analogique ;
- ▶ il est encore possible de raccorder un PABX traditionnel par un lien RNIS classique *via* une conversion réalisée chez l'opérateur ;
- ▶ Les IPBX pouvant être raccordés directement en IP au réseau de l'opérateur.

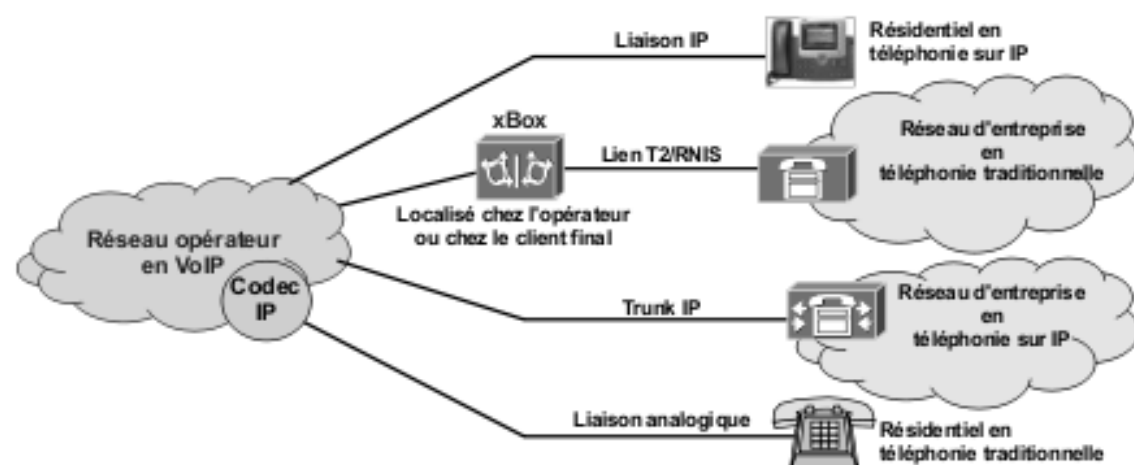


Figure 35.7 Le raccordement à un réseau voix sur IP.

11

La sécurité des systèmes d'informations

36

La sécurité des données

36.1 Généralités

L'ouverture des réseaux vers l'extérieur de l'entreprise et la multiplication des moyens d'accès fragilisent le système d'information. Ce dernier devient alors la cible d'attaques qui visent non seulement à prendre connaissance ou à modifier l'information, mais aussi à paralyser le système. Les moyens mis en œuvre pour protéger le réseau se regroupent sous le vocable de « sécurité des systèmes d'information ». Cependant, il convient de distinguer deux approches de la sécurité :

- ▶ la **sûreté de fonctionnement** (*safety*), qui concerne l'ensemble des mesures prises et des moyens utilisés pour se prémunir contre les dysfonctionnements du système ;
- ▶ la **sécurité** (*security*), proprement dite, qui regroupe tous les moyens et les mesures prises pour mettre le système d'information à l'abri de toute agression.

L'ouverture des réseaux de l'entreprise au monde extérieur, la décentralisation des traitements et des données ainsi que la multiplication des postes de travail accroissent les risques de dénaturation des systèmes et d'altération des données. Les différentes attaques peuvent donc s'exercer à tous les niveaux du système, une bonne sécurisation doit être pensée avec une approche structurée, non en termes de menaces à contrer mais en termes de services à protéger en se posant la question d'où provient et où peut s'exercer la menace. En fait, les menaces peuvent se regrouper en cinq catégories, celles qui visent à :

- ▶ prendre connaissance des données sans y être habilité (**confidentialité**) ;
- ▶ altérer les données (**intégrité**) ;

- ▶ mystifier les correspondants par usurpation d'identité (**authentification**) ;
- ▶ nier l'existence d'une transaction (**désaveu** ou répudiation) ;
- ▶ paralyser les systèmes (**déni de service**).

Les mécanismes mis en œuvre pour garantir la confidentialité, l'intégrité, l'authentification, le non-désaveu et la disponibilité du système peuvent se répartir en deux techniques : celles qui tendent à protéger les données et celles qui tendent à protéger les systèmes.

36.2 La protection des données

D'une manière générale, la confidentialité est assurée par le chiffrement des messages, l'authentification des correspondants par un échange de mots de passe plus ou moins simple, enfin le non-désaveu est garanti par un système d'accusé de réception ou par l'intervention d'un tiers (le notaire) qui mémorise et authentifie les transactions (notarisation).

36.2.1 Notions de cryptographie

■ Généralités

Le chiffrement est une technique destinée à rendre les données inintelligibles pour des tiers non autorisés. L'opération de brouillage du texte s'effectue à partir d'une clé (clé de chiffrement). La figure 36.1 illustre une chaîne de cryptage. Le message en clair est codé (chiffé) à l'aide d'une clé de chiffrement ; seul, le cryptogramme (message chiffré) est transmis sur le réseau. Le destinataire du message effectue le décryptage à l'aide d'une clé de déchiffrement.

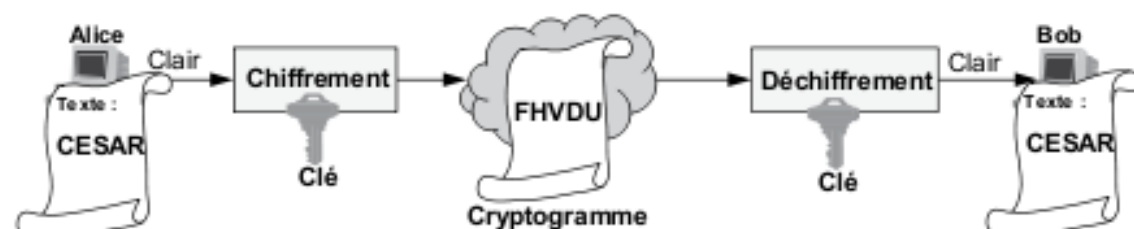


Figure 36.1 Principe de la cryptographie.

Les techniques de cryptographie sont utilisées pour :

- ▶ assurer la confidentialité des données (algorithme de chiffrement),
- ▶ garantir l'intégrité des données (algorithme de hachage),
- ▶ authentifier l'émetteur des données (algorithme de signature numérique).

■ Les méthodes de chiffrement symétrique

Les systèmes à **clés symétriques** ou secrètes utilisent une clé de chiffrement et une clé de déchiffrement identiques, convenues par avance et conservées secrètes (algorithme à clé secrète). Ils ne permettent pas d'identifier l'interlocuteur distant. Les algorithmes utilisent deux techniques : la substitution et la transposition indépendamment l'une de l'autre ou successivement.

Le code de César, dit aussi code à translation, est le plus vieux algorithme de chiffrement symétrique à substitution connu. Son principe est extrêmement simple, il suffit de substituer à chaque lettre du clair, une lettre de l'alphabet obtenue par simple translation (clé secrète) dans l'alphabet. Par exemple, si la translation est de 3, la lettre A est remplacée par la lettre D, la lettre B par E... La figure 36.2 illustre l'application de ce codage au mot « CESAR », la clé étant fixée à 4. Le clair « CESAR » donne alors pour chiffre le message « FHVDU ».

	A	B	C	D	E	F	G	H	I	J	K	L	M	N	O	P	Q	R	S	T	U	V	W	X	Y	Z
C			.	.	.	F																				
E					.	.	.	H																		
S																			.	.	.	V				
A	.	.	.	D																						
R																		.	.	.	U					

Figure 36.2 Le décalage de César.

Le **DES** (*Data Encryption Standard*) d'origine IBM (Karl Meyer 1977) est l'algorithme à clé symétrique le plus connu. Il consiste en une suite de substitutions (DES-S) et de transpositions, ou permutations (DES-P), par bloc de 64 bits. La figure 36.3 illustre de manière simple le principe d'un tel code. Utilisant une clé de 56 bits (64 bits dont 8 de parité), le DES est aujourd'hui facilement « cassable », il est remplacé par le triple DES (3DES, application de trois DES successivement avec trois clés indépendantes).

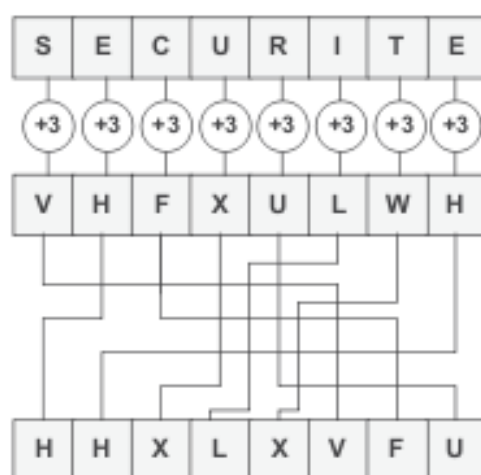


Figure 36.3 Principe du DES.

Les algorithmes de cryptographie à clé secrète demandent relativement peu de puissance, le temps de calcul est compatible avec un échange interactif de messages. Cependant, la découverte de la clé secrète donne accès à l'information. Dans de tels algorithmes, le secret (la clé) doit être transmis, d'où les risques d'interception, ou alors préalablement connu des deux correspondants.

■ Les méthodes de chiffrement asymétrique

Évitant la diffusion de clés, les systèmes à **clés asymétriques** utilisent deux clés, l'une est connue de tous (**clé publique**), l'autre n'est connue que de l'un des correspondants (**clé secrète**). Le message chiffré avec l'une ne peut être déchiffré qu'avec l'autre. Les deux clés sont reliées mathématiquement entre elles, mais l'utilisation de grands nombres rend ce lien pratiquement impossible à retrouver. La figure 36.4 illustre ce mécanisme.



Figure 36.4 Principe de la cryptographie à clé publique.

Le système de cryptographie à clé asymétrique le plus répandu, le **RSA** du nom de ses inventeurs (*Rivest, Shamir et Adleman*), repose sur l'arithmétique des grands nombres. Fondés sur la difficulté de factoriser des nombres premiers, les systèmes à clé publique permettent d'assurer la confidentialité des données mais, aussi d'authentifier l'émetteur d'un message.

■ L'authentification de l'émetteur

Un message chiffré avec la clé publique n'est déchiffrable qu'à l'aide de la clé privée, cela assure la confidentialité mais ne permet pas d'authentifier l'auteur du message. L'authentification de l'émetteur peut être obtenue en chiffrant le message avec la clé privée et en le déchiffrant avec la clé publique (figure 36.5).



Figure 36.5 Principe de l'authentification de l'émetteur.

Si Alice, à l'aide de la clé publique de Bob, déchiffre le message, c'est que celui-ci a bien été codé à l'aide de la clé privée de Bob, donc Bob est bien l'émetteur du message. Ce procédé ne garantit pas la confidentialité des messages, tout possesseur de la clé publique peut déchiffrer le message, il ne garantit que l'origine (le détenteur de la clé privée), c'est un système de signature de messages.

■ Le protocole d'échange de clés Diffie-Hellman

La cryptographie à clé publique nécessite une puissance de calcul importante. Le protocole d'échange de clés de Diffie-Hellman permet de construire une clé secrète (clé de session) sans que celle-ci circule sur le réseau. L'initiateur de l'échange transmet à son correspondant deux nombres grands et premiers (g , n). Les correspondants déterminent une clé privée, tenue secrète. Chacun, à partir de g , n et de sa clé secrète

(nombres aléatoires A et B), génère une clé publique et la communique à l'autre. Puis, chacun à partir de sa clé privée, de sa clé publique et de la clé publique de son correspondant, calcule la clé de session (figure 36.6).

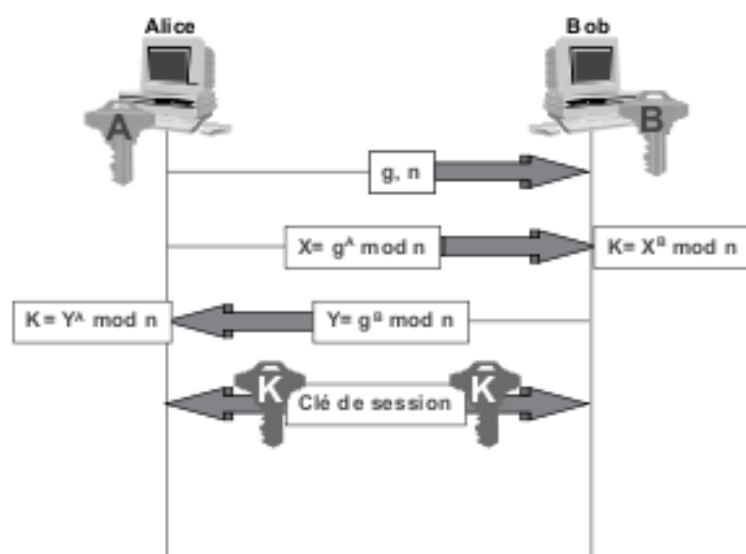


Figure 36.6 Principe de l'échange de Diffie-Hellman.

Le protocole de Diffie-Hellman permet de sécuriser l'échange de clés, cette technique est utilisée dans IPSec (*IP Secure*).

■ Contrôle de l'intégrité du message

Pour vérifier l'intégrité d'un message, on utilise une technique similaire à celle du CRC (*Cyclic Redundancy Check*). Une fonction dite de hachage (hash) est appliquée au contenu du message. Le résultat obtenu ou digest (résumé, sceau...) est joint au message à transmettre, il est recalculé par le destinataire. Si le résultat du calcul local est identique au *digest* reçu, le message n'a pas été altéré (figure 36.7).

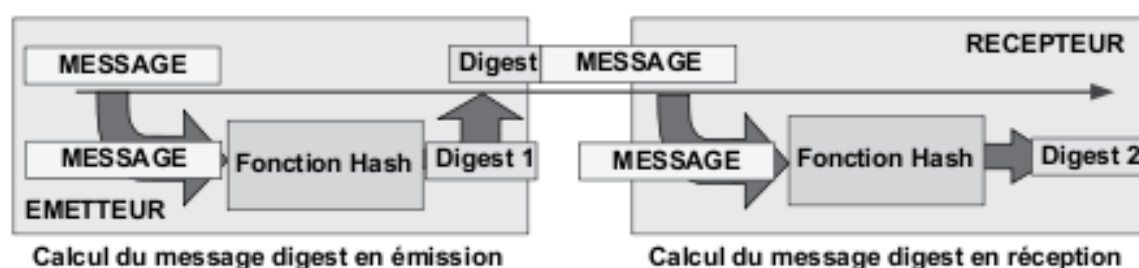


Figure 36.7 Principe de détermination du « digest ».

La fonction de hachage doit garantir qu'il est impossible à partir du *digest* de retrouver le message initial (non-retour arrière ou *one-way hash*) et qu'il doit être quasi impossible que deux messages différents donnent le même *digest* (résistance à la collision). Le *digest* a une longueur de 128 bits (MD2 à MD5, *Message Digest X*, défini par Ron Rivest et normalisé par le RFC 1321) ou de 160 bits (SHA-1, *Secure Hash Algorithm*).

■ La signature numérique d'un message

En combinant un système de cryptographie et une fonction de hachage, on peut à la fois garantir l'intégrité du message et son authentification (**MAC**, *Message Authentication Code*). Selon que l'on utilise un système de cryptographie à clé secrète ou publique, on obtient une signature numérique dite symétrique ou asymétrique.

37

La sécurisation des échanges

37.1 L'usurpation d'identité

L'un des problèmes de la cryptographie à clé publique est la possible intervention d'une tierce personne (figure 37.1). Lorsqu'Alice veut entrer en relation avec Bob en utilisant un système de cryptographie à clé publique, elle doit demander à Bob sa clé publique. Cet échange peut être intercepté par Charlie, un intrus malveillant, qui peut répondre en lieu et place de Bob avec sa propre clé publique. De cette manière, Charlie pourra se substituer à Bob lors des prochains échanges, Alice étant persuadée que les messages lui proviennent bien de Bob. La même opération est réalisée lorsqu'Alice envoie sa clé publique à Bob. Cette attaque est connue sous le nom de « *Man in the middle* ».



Figure 37.1 La substitution d'identité.

Afin d'éliminer la substitution d'identité, les clés publiques sont disponibles sur un serveur de clés publiques (annuaire) et donc accessibles à tous les utilisateurs, encore faut-il que soit confirmée la relation clé publique/

possesseur. C'est l'intervention d'un tiers de confiance (CA, *Certificate Authority*) qui garantit la correspondance entre une clé publique et son propriétaire par la délivrance d'un certificat. Le certificat contient l'identifiant d'un utilisateur et sa clé publique ; le certificat est signé avec la clé privée de l'autorité d'authentification. L'autorité de certification peut être interne à l'entreprise (disponible sur l'intranet de l'entreprise) ou être un prestataire de service de certification.

37.2 La sécurité et le protocole de transmission

37.2.1 La sécurité et le protocole PPP (RFC 1334)

Rappelons que le protocole PPP (*Point to Point Protocol*) assure quatre fonctionnalités (figure 37.2), la négociation des paramètres de connexion, l'affectation d'adresses IP, la sécurisation des échanges par authentification des communicants et enfin le transfert de données.

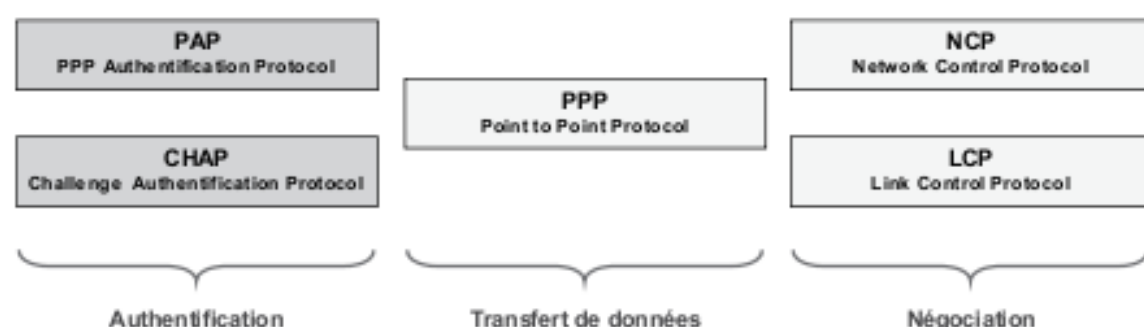


Figure 37.2 PPP et ses sous-protocoles.

PAP (*Password Authentication Protocol*) consiste en un simple échange de mots de passe en clair sur le réseau.

CHAP (*Challenge Handshake Authentication Protocol*) repose sur l'échange de messages cryptés selon une clé secrète (algorithme à clés symétriques) qui ne circule pas sur le réseau. L'identificateur envoie un

message en clair à l'interlocuteur distant (*Challenge*). Celui-ci crypte le message avec la clé secrète et le renvoie à l'identificateur (sceau). Si le message reçu est correctement crypté, l'identificateur en conclut que son interlocuteur est bien celui qu'il prétend être. La séquence exécutée dans les deux sens peut être répétée plusieurs fois au cours de la session (échange de sceaux).

37.2.2 Sécurisation des échanges sur le Web

■ S-HTTP (Secure HTTP)

S-HTTP introduit la cryptographie au niveau HTTP dont il constitue une extension. S-HTTP organise la session en trois étapes :

- ▶ l'authentification, par échange de mots de passe ;
- ▶ la négociation, phase où les interlocuteurs négocient le mode de cryptage à utiliser (DES, RSA...) ;
- ▶ la transaction, échange de messages cryptés selon le mode prédéfini.

La cryptographie est mise en œuvre par des scripts CGI ou par des « *daemon* » HTTP (*Plug in*).

■ SSL (Secure Sockets Layer)

Développé par Netscape Communication et intégré aux principaux navigateurs, SSL constitue une couche insérée entre la couche application et la couche TCP.

SSL crée une connexion qui permet un échange sécurisé dans le réseau. Cependant, le système n'est pas infallible. En effet, la clé secrète est générée à partir de l'horloge de la machine, ce qui permet de la trouver en quelques minutes.

■ TLS, DTLS

Le protocole TLS (*Transport Layer Security*, RFC 2246) est la version standardisée du protocole SSL. Bien que dérivé de SSL, TLS est incompatible avec ce dernier mais, si le correspondant ne met pas en œuvre TLS, il est capable de basculer en fonctionnement SSL. Il autorise l'authentification

par échange de certificats (X.509 v3), la confidentialité par chiffrement des données, la compression et la détection d'une éventuelle corruption des données. Son principal avantage vis-à-vis d'IPSec est qu'il ne nécessite aucune adaptation du réseau. Il utilise une connexion fiable (TCP) et sécurise SMTP, NNTP, http (HTTPS).

37.2.3 IP Security

Non limité aux échanges *via* un navigateur, l'*Internet Protocol Security Standard* fournit une sécurisation au niveau IP. Développé à l'origine le cadre d'IPv6 et adapté à IPv4, IPSec offre les services de contrôle d'accès, d'authentification, d'intégrité et de confidentialité des données, il met en œuvre un mécanisme d'anti-rejeu.

IPSec supporte de nombreux algorithmes de chiffrement (DES, triple DES, RC5, IDEA...), de hachage (MD5, SHA-1...) et d'authentification (signatures RSA ou DSS, clé secrète, clé publique). Dans ces conditions, l'utilisation d'IPSec est précédée d'une phase de négociation pour déterminer les mécanismes qui seront utilisés. L'ensemble des informations partagées entre les deux systèmes, pour établir une communication sécurisée, constitue une association de sécurité (SA, *Security Association*). Une association de sécurité est unidirectionnelle, un échange de données *full duplex* aboutit à la création de deux associations de sécurité.

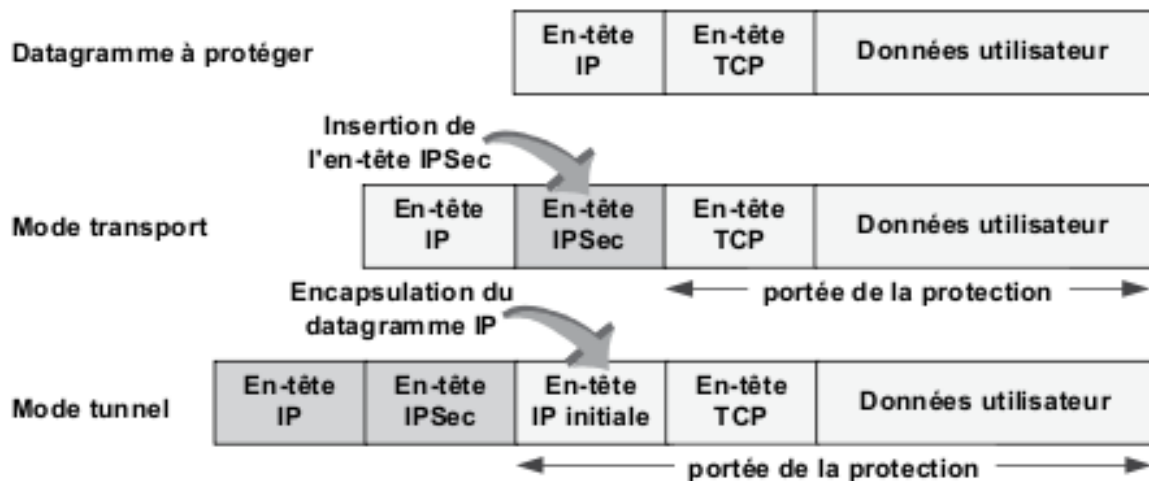


Figure 37.3 IPSec, mode transport et mode tunnel.

IPSec intègre deux modes de travail, le **mode transport** et le **mode tunnel** (figure 37.3). Le mode transport ne protège que le champ Données du datagramme IP. Le mode tunnel encapsule le datagramme IP d'origine, un nouvel en-tête IP est ainsi ajouté, les adresses IP source et destination initiales sont ainsi masquées.

38

La sécurisation du réseau

38.1 Les menaces

Les menaces contre les systèmes visent essentiellement à les rendre inaccessibles ou à en altérer profondément les performances. Par exemple, le protocole **ICMP** (*Internet Control Messages Protocol*) constitue l'une des failles (vulnérabilités) des environnements TCP/IP. En effet, il suffit, par exemple, d'adresser à un routeur d'interconnexion des paquets ICMP de signalisation de congestion pour que le routeur destinataire ralentisse ses émissions de messages vers le réseau extérieur.

De même, les paquets ICMP sont utilisés par le protocole RIP (ICMP *redirect*) pour modifier les tables de routage. Il est alors possible de faire croire à un routeur qu'il n'existe plus aucune route pour aller vers tel ou tel site, ou même détourner le trafic vers un autre site. Ces attaques sont généralement difficiles à détecter.

Les attaques peuvent se classer en deux catégories : celles qui visent à prendre connaissance d'informations pour les exploiter ou les altérer et celles qui visent à paralyser voire détruire les systèmes. Les modes d'attaque sont nombreux, ils vont de la simple usurpation de mots de passe (*brute force attack*¹ ou *dictionary attack*²...) à l'introduction de code malicieux (virus) en passant par la mystification des systèmes (*IP Spoofing*³...).

1 Tentative de pénétrer un système en essayant toutes les combinaisons possibles de mots de passe.

2 Ces attaques visent à déchiffrer les mots de passe encryptés par comparaison avec un dictionnaire de mots de passe chiffrés.

3 IP Spoofing consiste à modifier les adresses IP pour intercepter un trafic.

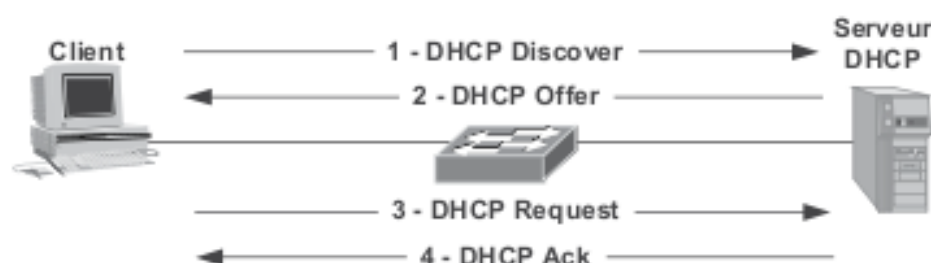
38.2 La protection de l'intranet

38.2.1 Protection du réseau local en interne

La sécurisation de l'infrastructure locale recouvre deux techniques. La première tente de prévenir les connexions non autorisées par le contrôle d'adresses (association d'un port du commutateur ou du *hub* et d'une adresse MAC) et la désactivation des ports non utilisés. La seconde assure un cloisonnement des trafics par la constitution de VLAN. Cependant, le service de DHCP reste le service le plus vulnérable sur une infrastructure LAN.

■ Exemple : sécurisation du DHCP (partielle)

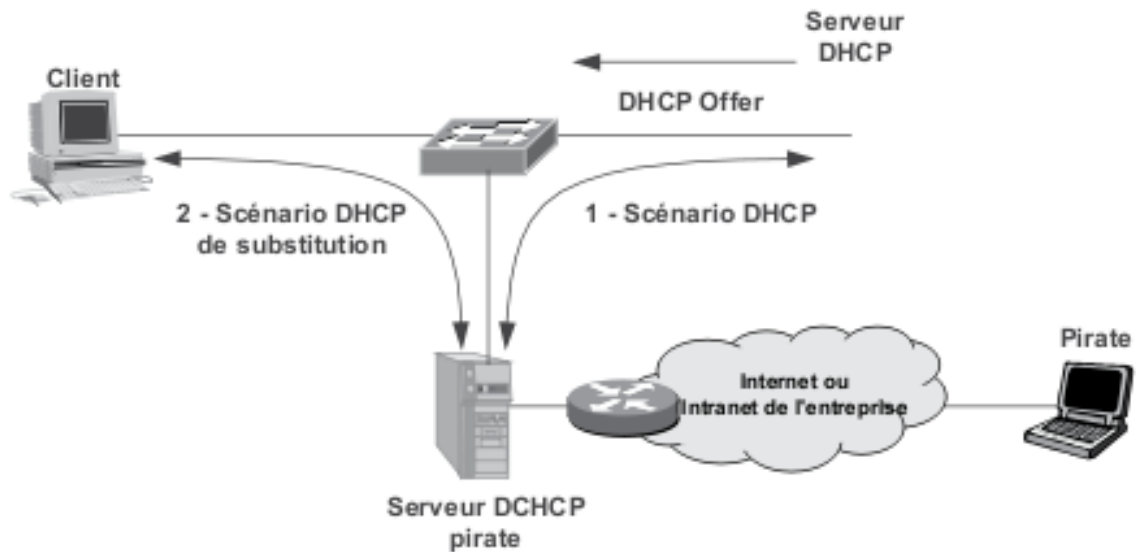
Rappelons que le DHCP (*Dynamic Host Configuration Protocol*) a pour finalité de permettre à une station vierge de toute information sur le réseau de se raccorder à celui-ci en obtenant du serveur DHCP tous les paramètres de configuration. La figure 38.1 rappelle le fonctionnement d'un service de DHCP.



- 1 - Message d'exploration (Broadcast), découverte des serveurs DHCP (DHCP Discover)
- 2 - Les serveurs DHCP actifs formulent une offre (DHCP Offer)
- 3 - Diffusion du serveur choisi (Broadcast), et demande de configuration (DHCP Request)
- 4 - Le serveur choisi acquitte et formule une offre.

Figure 38.1 Principe du service DHCP.

La vulnérabilité réside dans le fait que la requête est émise en mode *broadcast*, le serveur DHCP n'étant pas formellement désigné, un serveur « pirate » peut se substituer à lui et répondre, donnant alors à la station cliente une adresse valide mais par exemple une adresse de passerelle par défaut erronée permettant le détournement de trafic (figure 38.2).



- 1 - Le DHCP pirate, obtient une configuration, il apprend les paramètres réseau
- 2 - A une requête DHCP discover, il répond plus vite que le serveur officiel, et est choisi
- 3 - Il se déclare alors comme passerelle par défaut et relaie le trafic vers la vraie passerelle

Figure 38.2 Principe du détournement de service DHCP.

Plusieurs techniques peuvent être déployées pour sécuriser un service de DHCP, du simple DHCP *static* qui consiste à lier une adresse IP à une adresse MAC aux techniques les plus élaborées comme le DHCP *Relay* dans lequel le commutateur de rattachement ou la passerelle remplit une fonction de « NAT » en transformant le *broadcast* DHCP en *unicast*. Les requêtes DHCP sont identifiées par les ports sources (client DHCP UDP/68) et destination (serveur DHCP UDP/67).

38.2.2 Filtrage du trafic par le routeur d'accès

Le moyen le plus simple de protéger le réseau contre les intrusions peut être réalisé avec le routeur d'accès. Celui-ci assure des fonctions simples de filtrage par analyse des adresses source et destination. Le routeur n'a de visibilité que sur les données protocolaires du niveau 3, c'est-à-dire les adresses et, dans le mode IP, le protocole transporté dans le datagramme. Ses possibilités de filtre sont donc réduites à ces deux éléments, la sécurité offerte est faible. Les règles de filtrage sont réunies dans des listes (ACL, *Access Control List*). Le tableau 38.1 fournit un exemple de règles de filtrage, où « * » signifie toute valeur.

Tableau 38.1 Exemple de règles de filtrage.

Action	Protocole	Source	Destination	Commentaire
Accept	*	194.23.10.0/24	194.23.11.0/24	Trafic sortant vers l'établissement de Paris
Accept	*	194.23.11.0/24	194.23.10.0/24	Trafic entrant de l'établissement de Paris
Rejet	*	*	*	

Le filtre du tableau 38.1 n'autorise le trafic qu'entre deux établissements de l'entreprise. L'établissement de filtres est délicat, il nécessite une analyse fine des trafics autorisés et des trafics interdits. Dans notre exemple simple, l'écriture de la ligne 3 en tête de liste interdirait tout trafic.

38.2.3 La translation d'adresses (RFC 1631)

La translation d'adresses est un moyen de contourner la pénurie d'adresses Internet, mais aussi de masquer, vis-à-vis de l'extérieur, le plan d'adressage de l'entreprise (*IP Masquerade*).

La traduction statique fait correspondre à une adresse interne du réseau une adresse externe, généralement publique. Ce mode de translation résout à la fois le problème de la pénurie d'adresse, du masquage du plan d'adressage local (*masquerade*) et sécurise le réseau en n'autorisant que certaines stations à accéder à l'Internet.

La translation statique limite le nombre de machines ayant accès à l'extérieur au nombre d'adresses publiques attribuées. La traduction dynamique s'affranchit de cette limite. Lorsqu'une machine veut atteindre une machine extérieure, le NAT (*Network Address Translation*) associe à l'adresse locale interne une adresse globale interne, ou adresse externe, choisie parmi un pool d'adresses mises à sa disposition. Le NAT introduit un protocole à état, indépendamment du fait qu'en cas de défaillance du NAT les relations sont perdues, l'état créé doit être détruit en fin de communication et l'adresse attribuée rendue disponible pour une autre connexion vers l'extérieur. Un temporisateur est donc associé à chaque connexion, il est réinitialisé à chaque message, la connexion est libérée sur *time out*.

Cependant, le nombre d'adresses publiques attribuées peut être insuffisant. Le **NAPT** (*Network Address and Port Translation*) permet à plusieurs machines de partager une même adresse externe par translation du numéro de port (figure 38.3). La fonction dite du **PAT** (*Port Address Translation*) autorise plusieurs milliers de connexions (65535) à se partager une même adresse IP externe dite aussi globale interne.

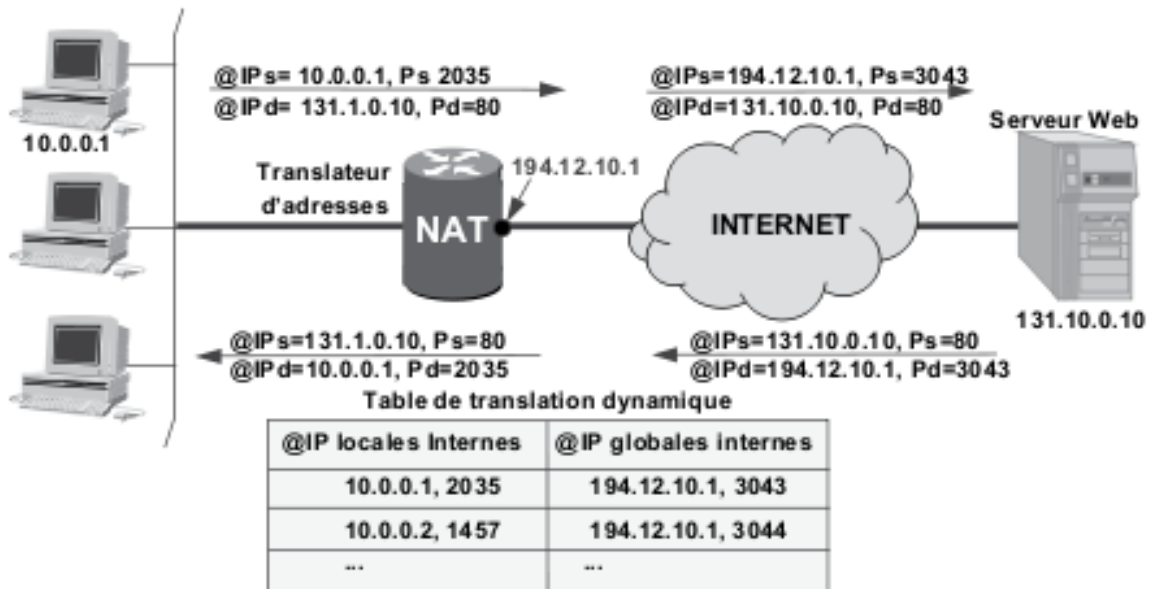


Figure 38.3 Principe de la translation de port.

En translation dynamique, la correspondance @IP interne/@IP externe est initialisée par la machine interne. Un datagramme entrant, sans correspondance dans la table de translation, est rejeté. Pour donner accès aux machines extérieures, à certains services, on peut utiliser la technique dite du « *Port forwarding* ». Celle-ci permet un accès direct au service concerné et un accès indirect *via* les services d'un **DNS** (*Domain Name System*).

38.2.4 Les pare-feu (*firewall*)

Le pare-feu (*firewall*) est un système aux fonctions de filtrage évoluées. Indépendamment des fonctions de routage et de translation d'adresses,

chaque paquet reçu est examiné, une décision de rejet ou d'acceptation est prise en fonction de nombreux critères :

- ▶ l'adresse destination,
- ▶ l'adresse source,
- ▶ le protocole transporté (ICMP, UDP...),
- ▶ le port destination,
- ▶ le port source,
- ▶ la valeur de certains flags (ACK, SYN...)...

On distingue trois types de pare-feu :

- ▶ ceux qui examinent le trafic et contrôlent celui-ci par simple filtrage selon des règles prédéfinies, ces pare-feu ne connaissent que les datagrammes qu'ils examinent, ils ignorent le contexte applicatif qui les a générés, ils sont dits pare-feu sans état ou *firewall stateless* ;
- ▶ les deuxièmes assurent le suivi des connexions établies, ils exercent leur contrôle dans un certain contexte, ils maintiennent un état pour suivre l'évolution d'une connexion, ils sont dits pare-feu à état ou *firewall statefull* ;
- ▶ enfin, les troisièmes exercent leur contrôle au niveau de l'application, ce sont les pare-feu applicatifs. Ils sont vus par les agents clients comme étant le serveur applicatif, ils sont généralement appelés « *Proxy server* ».

La figure 38.4 illustre les différentes architectures de sécurité envisageables, la mise à disposition d'un serveur public (service web, messagerie...) est généralement réalisée par la constitution d'une zone de sécurité dite **DMZ** (*DeMilitarized Zone*). Différentes zones de sécurité peuvent être constituées, chacune accessible selon des critères spécifiques (filtres). La zone démilitarisée accueillera les différents serveurs accessibles à la fois par le personnel de l'entreprise et par le monde extérieur. Pour différencier les services offerts et les règles de filtrage, il est possible de définir plusieurs DMZ, dans ce cas généralement l'une est accessible à tous (DMZ publique), et l'autre aux personnels de l'entreprise (DMZ privée).

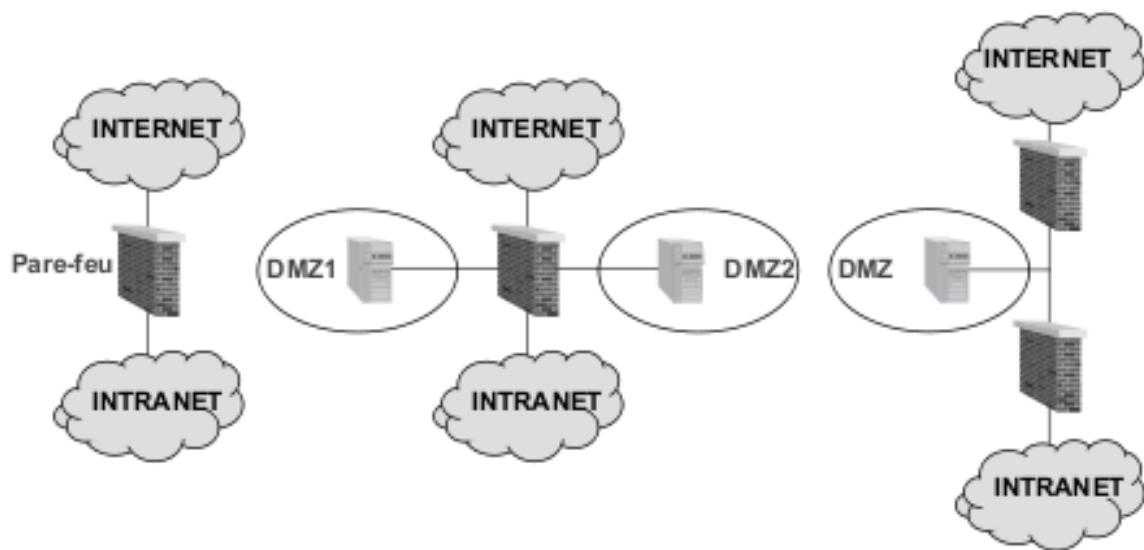


Figure 38.4 Les différentes architectures de sécurité.

La définition des filtres est similaire à celle réalisée pour les routeurs, seule, la portée de l'analyse est plus profonde. Tout datagramme non autorisé est rejeté. En cas de tentative de violation d'une règle, les pare-feu émettent des alertes. Le tableau 38.2 présente quelques exemples de règles de filtrage.

Tableau 38.2 Les exemples de règles de filtrage.

Règles	Destination		Source		Flag	Action	Commentaire
	Adresse	Port	Adresse	Port			
1	Externe	25	Interne	> 1023		accept	Connexion vers serveurs SMTP
2	Interne	> 1023	Externe	25	ACK	accept	Réponses aux connexions SMTP
3	Externe	23	Interne	> 1023		accept	Connexion à des services Telnet
4	10.0.0.5	23	194.28.12.1	> 1023		accept	Station externe autorisée Telnet
5	194.28.12.1	> 1023	10.0.0.5	23		accept	Trafic Telnet station 194.28.12.1
6	Interne	23	Externe	> 1023	ACK	accept	Réponses au trafic Telnet

L'énoncé des règles 1 et 2 autorise le trafic de toutes les stations du réseau interne vers des serveurs SMTP externes, mais n'autorise pas les connexions d'origine externe en provenance d'un service sur le port 25 puisque les messages d'origine externe doivent avoir le bit ACK (TCP) positionné, de même pour les règles 2 et 6. Cependant, pour autoriser une station spécifique à ouvrir depuis l'extérieur une session Telnet, nous avons inséré les règles 4 et 5. Si celles-ci avaient été placées après la règle 6, aucune connexion en provenance de l'extérieur à destination d'un service Telnet interne n'aurait été autorisée (bit ACK).

38.2.5 Les codes malicieux (virus)

Un virus est un programme « parasite » qui s'attache à un programme principal dont il modifie l'environnement de travail avec un objectif généralement destructeur (vers, chevaux de Troie, bombes logiques...). Les programmes virus ont aussi la possibilité de se propager de machine en machine directement avec le programme infecté (copie de programme), mais aujourd'hui de plus en plus par exploitation du carnet d'adresses de la machine infectée.

Les virus fonctionnent en tâche fond. Lorsqu'une certaine condition (date, type d'activité...) est réalisée, le virus effectue la tâche pour laquelle il a été programmé. Les virus peuvent altérer les données, les diffuser vers des adresses aléatoires ou pré-programmées, modifier le comportement du système allant de l'instabilité à la paralysie, voire à la destruction de certains composants du système (effacement du BIOS...). Aujourd'hui, les codes malicieux visent plus à obtenir des gains financiers qu'à détruire les systèmes.

Des logiciels, dits antivirus, permettent de se protéger des virus connus. Cependant, malgré les mises à jour, les « pirates » ont toujours un virus d'avance. La seule parade efficace consiste à n'échanger des données avec personne et à ne jamais raccorder son ordinateur à un réseau !

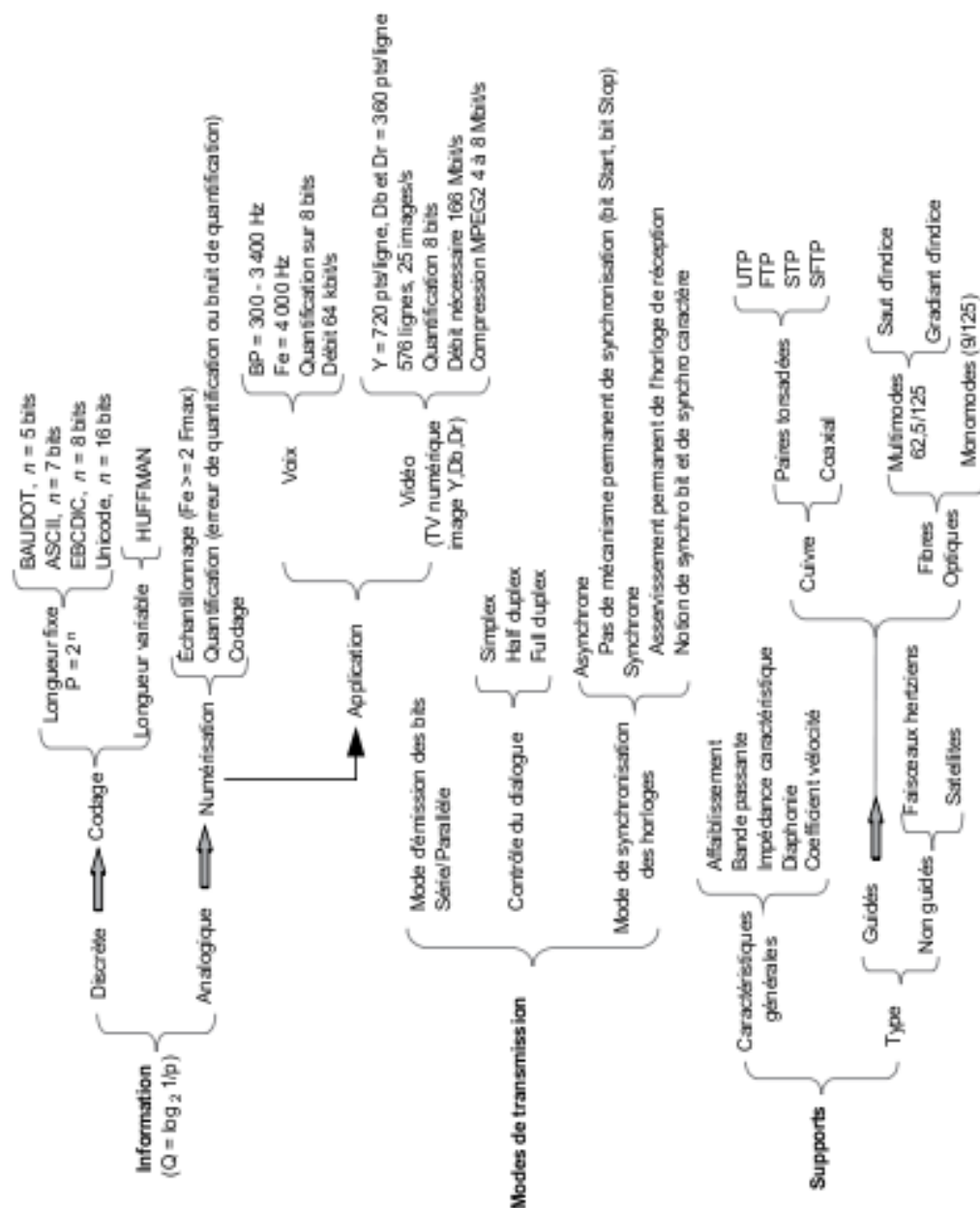
38.3 Conclusion

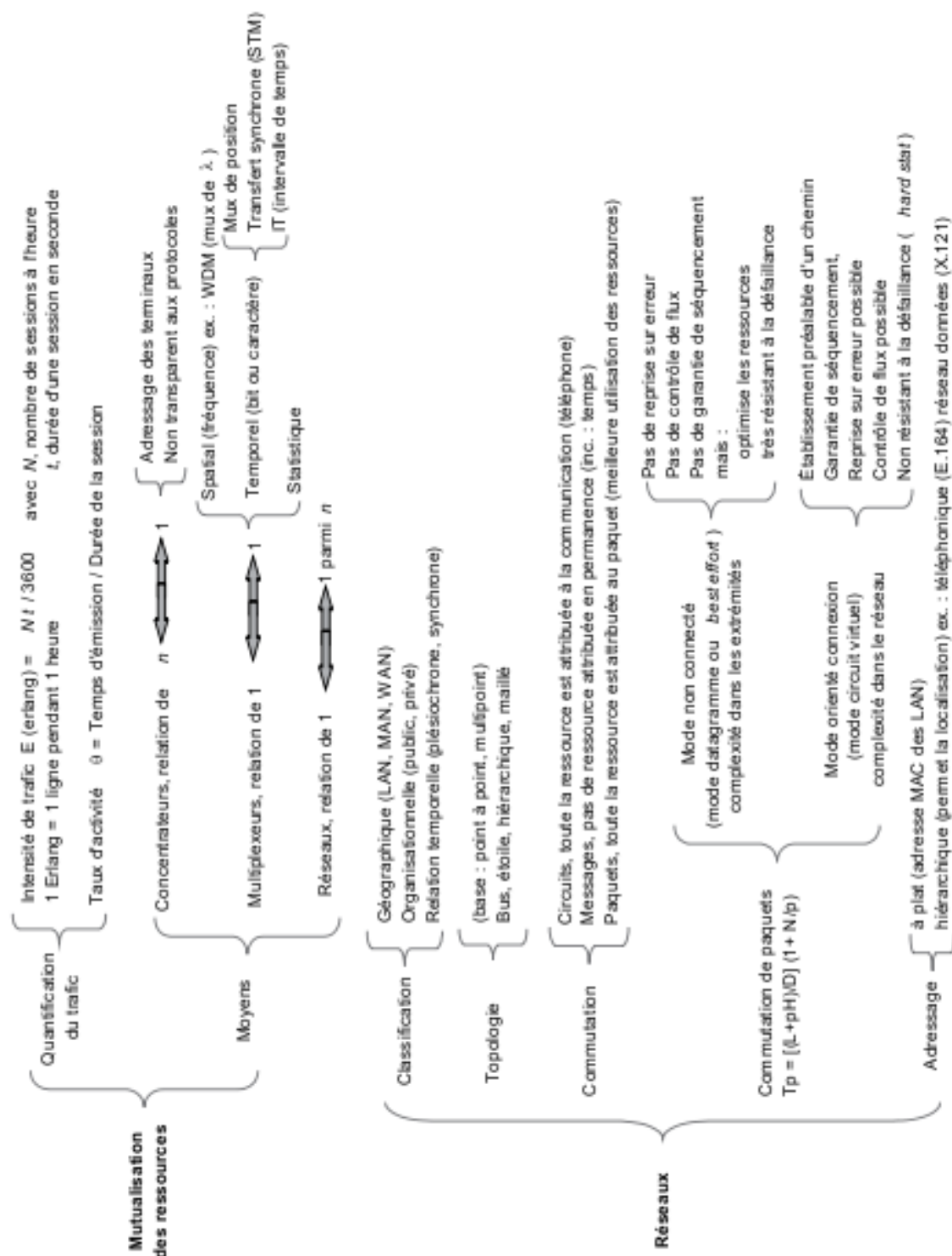
Quels que soient les moyens mis en œuvre, du fait de l'ouverture vers l'extérieur, les réseaux seront de plus en plus vulnérables. Dans ses conditions, il appartient à chaque entreprise de mesurer le risque et le coût d'une indisponibilité du système, de la divulgation d'information... et déterminer une politique dite de sécurité, d'en mesurer les coûts et d'assurer en permanence une veille technologique pour que les moyens employés restent efficaces devant l'évolution des menaces.

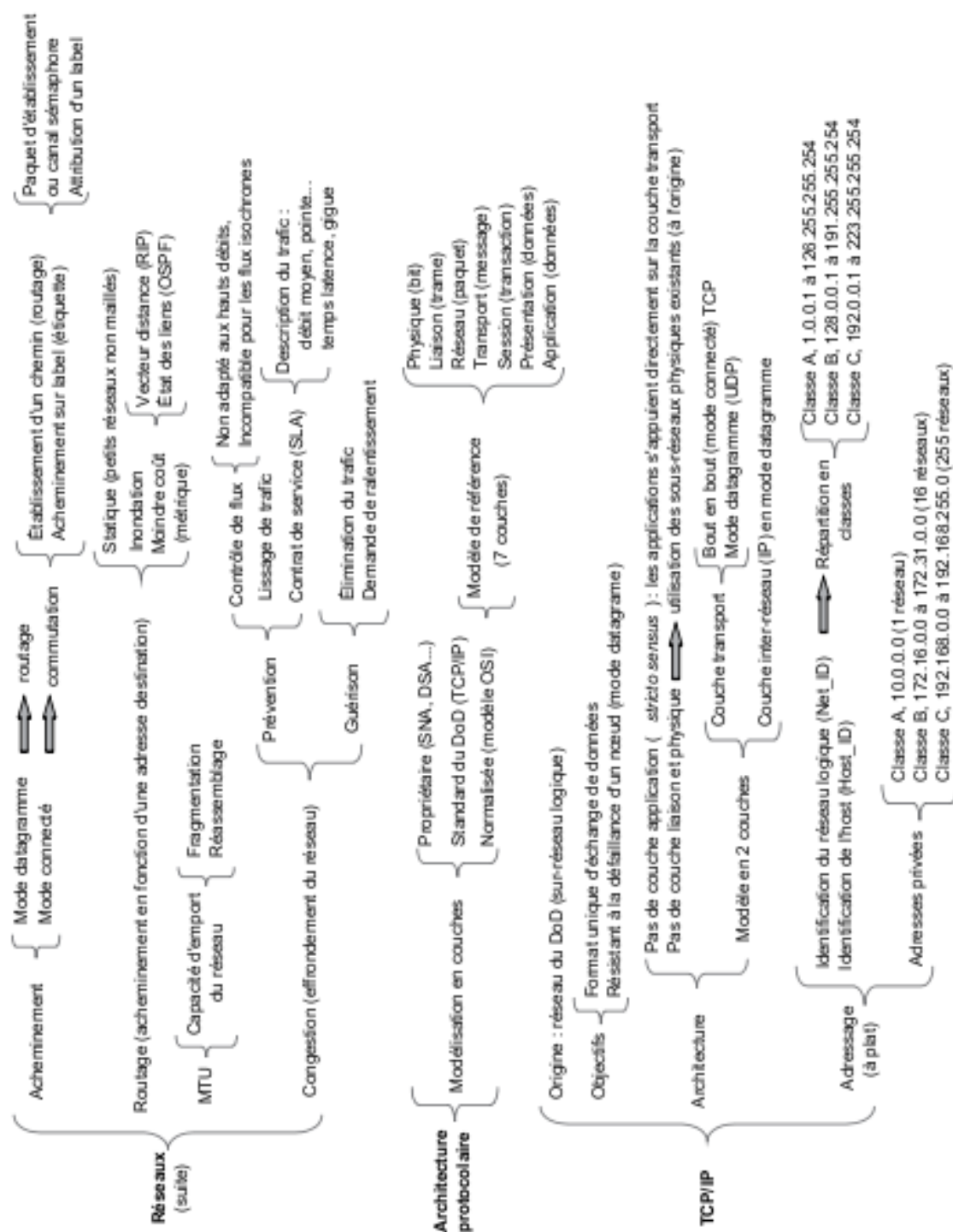
Annexes

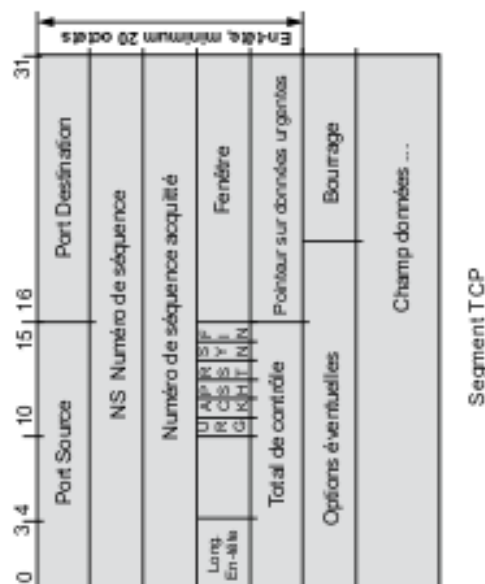
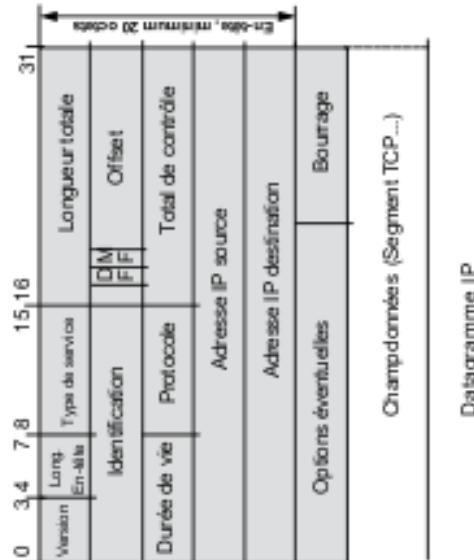
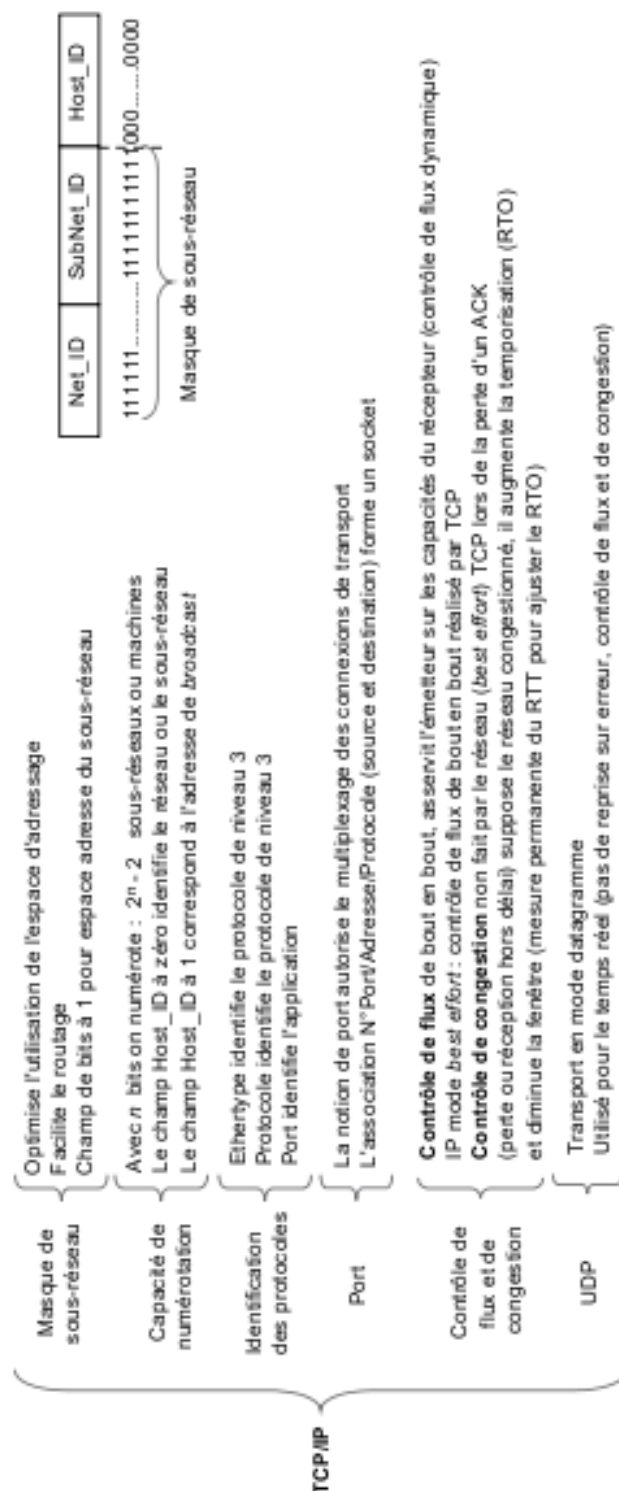
1) Synthèse

2) Normalisation









2

Normalisation

La normalisation peut être vue comme un ensemble de règles destinées à satisfaire un besoin de manière similaire. La normalisation dans un domaine technique assure une réduction des coûts d'étude, la rationalisation de la fabrication et garantit un marché plus vaste. Pour le consommateur, la normalisation est une garantie d'interfonctionnement, d'indépendance vis-à-vis d'un fournisseur et de pérennité des investissements.

En matière de télécommunication, la normalisation est issue d'organismes divers. Du groupement de constructeurs aux organismes internationaux, la normalisation couvre tous les domaines de la communication. D'une manière générale, elle ne s'impose pas, sauf celle émanant de l'**ETSI** (*European Telecommunications Standard Institute*) qui normalise les réseaux publics et leurs moyens d'accès.

Les principaux groupements de constructeurs sont :

ECMA (*European Computer Manufactures Association*), à l'origine constituée uniquement de constructeurs européens (Bull, Philips, Siemens...), l'ECMA comprend aujourd'hui tous les grands constructeurs mondiaux (DEC, IBM, NEC, Unisys...). En matière de télécommunications, l'ECMA comprend deux comités : le TC23 pour l'interconnexion des systèmes ouverts et le TC24 pour les protocoles de communication ;

EIA (*Electronic Industries Association*) connue, essentiellement, pour les recommandations RS232C, 449 et 442.

Les principaux organismes nationaux auxquels participent des industriels, administrations et utilisateurs sont :

AFNOR, Association française de normalisation,

ANSI, American National Standard Institute (États-Unis),

DIN, *Deutsches Institut für Normung* (Allemagne), bien connu pour sa normalisation des connecteurs (prises DIN) ;

BSI, British Standard Institute (Grande-Bretagne).

Les organismes internationaux :

- ▶ **ISO**, *International Standardization Organization*, regroupe environ 90 pays. L'ISO est organisée en *Technical Committee* (TC) environ 200, divisés en *Sub-Committee* (SC) eux-mêmes subdivisés en *Working Group* (WG) ; la France y est représentée par l'AFNOR ;
- ▶ **CEI**, Commission électrotechnique internationale, affiliée à l'ISO en est la branche électricité ;
- ▶ **UIT-T**, Union internationale des télécommunications secteur des télécommunications, qui a succédé en 1996 au CCITT (Comité consultatif international télégraphie et téléphonie), publie des recommandations. Celles-ci sont éditées tous les 4 ans sous forme de recueils. Les domaines d'application sont identifiés par une lettre :
 - V, concerne les modems et les interfaces,
 - T, s'applique aux applications télématiques,
 - X, désigne les réseaux de transmission de données,
 - I, se rapporte au RNIS,
 - Q, intéresse la téléphonie et la signalisation.

L'**IEEE**, *Institute of Electrical and Electronics Engineers*, société savante constituée d'industriels et d'universitaires, est essentiellement connue par ses spécifications sur les bus d'instrumentation (IEEE 488) et par ses publications concernant les réseaux locaux (IEEE 802), reprises par l'ISO (IS 8802).

Le panorama serait incomplet si on omettait de citer l'**IAB**, *Internet Architecture Board*, qui a la charge de définir la politique à long terme d'Internet, tandis que l'**IETF** (*Internet Engineering Task Force*) assure par ses publications (**RFC**, *Request For Comments*) l'homogénéité de la communauté TCP/IP et Internet.

La rédaction d'une norme est une succession de publications, la durée entre le projet et la publication définitive peut être très longue. En effet, chaque partie tente d'y défendre ses intérêts économiques et commerciaux. D'une manière générale, un projet de normalisation est formalisé dans un document brouillon qui expose les concepts en cours de développement (*Draft*) ; lorsque ce document arrive à une forme stable, les « drafts » sont publiés (*Draft proposable*), chaque pays émet son avis (vote). Enfin, une forme quasi définitive est publiée, elle constitue une base de travail pour les constructeurs (*Draft International Standard*). La norme appelée *International Standard (IS)* est ensuite publiée.

Index

A

AAL 267
ACCM 91
ACR 25
adresse
 hiérarchique 116
 MAC 216
 privée 131
ADSL 280, 282
aire de routage 309
alternat 45
annuaire 192
anycast 145
ARCEP 278
area backbone 310
ARP 14, 181
AS 309
Assured Forwarding 313
ATA 337
ATM 261
authentification 366

B

BAL 201
bande de base 50
BAS 96
best effort 313

bit

 de bourrage 69
 de parité 70
 de start 48
 de stop 48
Bit de bourrage 82
BOOTP 185
boucle locale 278
brasseurs 264
broadcast 116, 217

C

Call Manager 339, 342
CATV 33
câble coaxial 33
Centrex 352
CHAP 93, 373
chiffrement
 asymétrique 368
 symétrique 367
CIDR 135
CIR 260
classe de services 269
classe d'adressage 128
clé
 asymétrique 368
 publique 368
 secrète 368

CLP 265
codage
 en ligne 51
 Manchester 52
collision 215
commutation
 de circuits 104
 de messages 105
 de paquets 106
confidentialité 365
connexion de transport 151
contrôle
 de congestion 165
 de flux 80, 165
 de flux dynamique 82
 de flux explicite 166
 de flux implicite 81
CRC 71
cryptographie 366
CVC 110
CVP 110
CWDN 62

D

datagramme
 IP 136
 IPv6 146
DCE (Data Circuit Equipment) 5
dégrouper
 partiel 279
 total 278
démarrage lent 167
dénier de service 366
Dense WDM 61
DES 367

désaveu 366
détection d'erreur
 par clé calculée 71
DHCP 185, 378
Diaphonie 25
Differentiated Services 312
Diffie-Hellman 369
DiffServ 286, 312
diffusion
 dirigée 131
 générale 130
DMT 280
DMZ 382
DNS 14, 193
DTE (Data Terminal Equipment) 5
DVMRP 317
DWDM 61, 62

E

E.164 346
ECN 138
EFLEX 26
EGP 304, 309
EIR 260
E-LAN 286
E-Ligne 286
eMail 201
encapsulation 298
End-Span 29
état des liens 308
ETCD 5
Ethernet Carrier Grade 285
Ethernet full duplex 231
EtherPhone 337
EtherType 16

E-Tree 286
 ETTD 5
 EVC 285
 Expedited Forwarding 313

F

faisceaux hertziens 42
 fanion 67
 FCS (Frame Check Sequence) 83
 FDM 60
 FEC 272
 fibre
 à gradient d'indice 37
 optique 34
 à saut d'indice 36
 firewall 381
 statefull 382
 stateless 382
 FLIB 274
 Fourier 19
 FQDN 196
 fragmentation 140
 Frame Relay 258
 fréquence de coupure 21
 FTP 15, 24, 32, 199
 anonyme 200
 full duplex 45

G

G.711 327
 G.721 327
 G.722 327
 G.722-1 328
 G.723 328
 G.723-1 328

G.726 328
 G.727 328
 G.728 328
 G.729 328
 G.729 a 328
 gatekeeper 343
 gestionnaire d'appels 339
 G.Lite 282

H

H.225 343
 H.245 343
 H.323 343
 half duplex 45
 hardphones 337
 HDLC (High Level Data Link
 Control) 74
 HDSL 283
 HEC 266
 Host_ID 128
 HTML 206
 HTTP 15, 206, 346

I

ICANN 194
 ICMP 15, 177, 377
 ICMPv4 177
 ICMPv6 178
 IDSL 283
 IEEE 802.1p/Q 234
 IGMP 316
 IGP 304
 IGRP 308
 IMAP4 203
 impédance caractéristique 22

intégrité 365
intervalle de temps 62
IP 12
IP multicast 315
IPPhone 337
IP Precedence 137, 311
ISDN 324
ISN 154
IT 62

L

LAN 3
largeur de bande 21
LASER 35
LCP 92
LED 35
liaison
 de données 5
 full duplex 45
 half duplex 45
 à l'alternat 45
 simplex 45
LLC 218
LLC1 218
Location Server 347
LRC 70
LSP 272
LSR 271

M

MAC 215
MAN 4
MAQ 58
masque de sous-réseau 134
MCU 345
MDA 203

Media Gateway 338
métrique 123
MHS 202
MIC 326
middleware 204
Mid-Span 29
MLPP 358
mode
 datagramme 109
 non connecté 109
modem 56
modulation 57
 de fréquence 57
 de phase 57
 d'amplitude 57
MOS 327
MPLS 113, 270
MSS 163
MTA 202
MTS 202
MTU 118
MUA 202
multicast 116, 129, 145, 217
multiplexage
 fréquentiel 60
 temporel 60
multiplexeur 59
MUX 59

N

NAPT 381
NAT 132, 380
NCP 94
Net_ID 128
NEXT 25
NNI 264

nommage 117
 NRZ (No Return to Zero) 52
 NVP 26
 Nyquist 55

O

OCH 39
 OEM 40
 Olympic Service 313
 OMS 39
 ondes électromagnétiques 40
 OSI 8
 OSPF 15, 308
 OTN 39
 OTS 39
 OUI 216

P

PABX 331
 PAP 93, 373
 Paradiaphonie 25
 pare-feu 381
 passerelle voix 338
 PAT 381
 piggybacking, 85
 PIM 318
 PING 178
 PoE 29
 ponts 300
 PoP 252
 POP3 203
 port 16
 bien connu 153
 référéncé 153
 PPP 15, 88

PPPoA 96, 282
 PPPoE 96, 282
 Premium Service 313
 protocole 16
 à anticipation 77
 Proxy server 348, 382
 PSFLEXT 26
 PSNEXT 26
 PSTN 323
 PSTN (Public Switched Telephone
 Network) 330, 337
 PVC 110

Q

Q.931 343
 QAM (Quadrature Amplitude
 Modulation) 58
 QoS 359

R

rapidité de modulation 53
 RARP 15, 183
 RAS 343
 Real Time Protocol 341
 RED 169
 redirecteur 213
 Redirect server 348
 Registrar 347, 348
 rejet
 sélectif 80
 simple 80
 réseau à commutation 101, 103
 résolution d'adresses 181
 Return Loss 25
 RFC 1356 299

RIP 15, 307
routage
 aire de 309
 à état des liens 125
 inter-domaine 304
 intra-domaine 304
 par inondation 123
 statique 121
 vecteur-distance 125, 307
routeur 300
 désigné 310
RPC 205
RSA 369
RTC 323
RTO 77, 157, 160
RTP 341
RTSP 347

S

SAP 15, 347
SC 39
SCCP 347
SDP 347
SDSL 282
Send and Wait 76
Service Level Agreement 259
services Platinum, Gold, Silver
 ou Bronze 138
Shannon 55, 325
S-HTTP 374
SIP 346
SLA 259, 313
SLIP 15, 73
SNAP 219
SNMP 15

socket 152, 206
softphones 337
sous-réseau 133
sûreté de fonctionnement 365
SSL 374
ST 39
STP 24
survivabilité 354
SVC 110
Switched Ethernet 230
synchronisation
 caractère 68
système autonome 309

T

T.120 344
table
 de routage 121
 d'acheminement 121
TCP 12, 151
TCP/ECN 138
TCP/IP 12, 14
TELNET 15
TFTP 15, 198
TLD 194
ToIP 336
TPKP 344
Traceroute 180
translation d'adresses 380
transmission analogique 57
transparence
 au caractère 68
 binaire 68
tunneling 298

U

UAC 347
UADSL 282
UDP 171
U-DWDM 62
UNI 264
unicast 115, 145, 217
URL 206
UTP 23, 32

V

valence 55
VCI 263
VDSL 283
VLAN 232
VoIP 336
voix paquetisée 336
VPI 264
VPN 275
VRRP 291

W

WAN 4
WDM 61, 62
Well known ports 153
WWW 206

X

X.25 255
xDSL 95
XON-XOFF 81

Z

zone 309